

Giovanni Organtini

FISICA SPERIMENTALE



VISIT...

LANZAROTE
Caliente.COM

FISICA SPERIMENTALE – VERS. 5 MAGGIO 2015
 © 2013–2015 Giovanni Organtini, Sapienza–
 Università di Roma & INFN–Sez. di Roma



Questo è un libro di testo elettronico gratuito e aperto. Questo testo è stato realizzato usando per quanto possibile strumenti *Open Source* ed è disponibile gratuitamente secondo i termini della Licenza Creative Commons Attribution–NoDerivs 3.0. Puoi copiarlo, stamparlo e ri-distribuirlo come vuoi. Non puoi realizzare opere derivate senza il consenso dell'autore.

Per maggiori informazioni sulle Licenze Creative Commons, consulta il sito <http://www.creativecommons.org>.

Il lavoro per la redazione di questo prodotto è in corso ed è solo all'inizio. Se la porzione di libro disponibile **ti è piaciuta**, per continuare a usufruire di questi servizi e permettere ad altri di farlo, considera la possibilità di donare una cifra a piacere, a partire da 50 centesimi, visitando il sito di [PayPal](#) specificando l'indirizzo e-mail Giovanni.Organtini@roma1.infn.it. Le donazioni consentiranno un costante aggiornamento dei contenuti.

Tutti i contributori, indipendentemente dall'importo donato e salvo esplicita indicazione contraria, saranno inseriti in una *mailing list* attraverso la quale verranno informati dell'uscita di nuove edizioni del volume. Donando per questo libro si accetta implicitamente l'inserimento nella lista e si dichiara il consenso al trattamento dei propri dati personali (nome, cognome e indirizzo di posta elettronica) ai sensi del DL 196/2003 e successive modificazioni. I dati personali verranno utilizzati al solo fine di informare i donatori dell'uscita di nuove edizioni del volume e verranno trattati nel rispetto della normativa. I contributori hanno la facoltà di chiedere di non essere inseriti in tale lista o di esserne rimossi, ai sensi dl DL sopra citato. È sufficiente indicare la volontà di esercitare questo diritto in sede

di versamento della donazione o inviando una mail all'autore del volume.

Per contributori s'intendono sia coloro che abbiano versato un contributo in denaro, sia coloro che avranno segnalato eventuali errori o avanzato proposte di riformulazione di parti del libro accolte dall'autore.

Puoi contattare l'autore di questo libro elettronico ai seguenti recapiti:

Prof. Giovanni Organtini
 "Sapienza", Università di Roma
 Dip.to di Fisica
 P.le Aldo Moro, 2
 00185 ROMA (Italy)

Tel: +39 06 4991 4329 Fax: +39 06 4453 829
 e-mail: giovanni.organtini@uniroma1.it

Risorse utili

Sono numerose le risorse a cui puoi attingere per rimanere aggiornato sulla fisica e sugli sviluppi di quest'opera. In particolare ti consigliamo di visitare periodicamente il sito di **FISICAST** all'indirizzo <http://www.fisicast.it>: si tratta di un *podcast* sul quale pubblichiamo, con cadenza mensile, brevi brani audio che parlano di fisica con un linguaggio semplice e accessibile a tutti, o quasi. Trovi i podcast di FISICAST anche su **iTunes** e su molti altri servizi di *podcasting*.

Puoi anche sottoscrivere il canale **twitter** dell'autore all'indirizzo <https://twitter.com/organtini> sul quale sono pubblicate notizie relative a quest'opera e su altri aspetti della fisica e non solo.

Se hai un account **Google+** puoi seguire i **post** dell'autore all'indirizzo <http://plus.google.com/+GiovanniOrgantini>.

Infine, puoi visitare il sito dell'autore all'indirizzo <http://www.roma1.infn.it/people/organtini> per attingere a ulteriori informazioni.

Ricorda che l'autore è disponibile ad accogliere suggerimenti o a esaudire richieste di integrazioni, nonché a tenere corsi o seminari per studenti, insegnanti o per un pubblico più generico. Se trovi errori (il che è sempre possibile) ti preghiamo di segnalarli. Ci aiuteranno a offrire un servizio migliore.

Indice

Prefazione	1	3.2.1 La regressione lineare	42
Alla scoperta dell'Universo	7	3.2.2 La costruzione di un modello	44
1 Scoprire la Fisica	11	3.3 La Temperatura	46
1.1 Il problema della misura	11	4 Calore e temperatura	51
1.2 La luce	11	4.1 La teoria del calorico	51
1.3 Le onde	11	4.2 Trasporto del calore	53
1.4 Il moto, il lavoro e l'energia	12	4.3 Falsificare una teoria	56
1.5 La termodinamica	12	5 Ottica geometrica	59
1.6 Campi di forze	13	5.1 Riflessione della luce	59
1.7 Correnti elettriche	13	5.2 Una prima interpretazione	64
1.8 Campi magnetici	13	5.3 La rifrazione	66
1.9 Corpi rigidi	13	5.4 Conferma della teoria corpuscolare .	68
1.10 Onde elettromagnetiche	13	5.5 Applicazioni	69
1.11 Ancora sulla natura della luce	13	6 Le onde e i fenomeni ondulatori	73
1.12 L'entropia e il secondo principio . . .	14	6.1 Caratterizzazione delle onde	74
1.13 Le particelle elementari	14	6.2 Riflessione e rifrazione delle onde . .	77
2 Il ruolo della misura in Fisica	15	7 Interferenza	79
2.1 Misure e teorie	16	7.1 Casi particolari	80
2.2 Le misure di base	18	7.2 Onde stazionarie	82
2.3 Gli strumenti	21	8 Effetti del moto sulle onde	85
2.4 Notazione scientifica	23	8.1 L'effetto Doppler	85
2.5 Un esperimento istruttivo	24	8.2 Effetto Doppler Relativistico	87
2.6 Proprietà statistiche delle variabili casuali	25	8.3 L'effetto Čerenkov	88
2.7 L'interpretazione delle misure	28	8.4 Il <i>red shift</i> delle galassie	90
2.8 Analisi statistica delle misure	30	9 La diffrazione	93
2.9 Errori sistematici	32	9.1 Sperimentiamo la diffrazione	94
2.10 Propagazione degli errori	34	9.2 Definiamo la natura della luce	95
2.10.1 La media pesata	37	9.3 La matematica della diffrazione . . .	97
3 Definire le grandezze fisiche	39	9.3.1 Diffrazione da fenditura	97
3.1 Massa e Peso	39		
3.2 La radioattività	41		

9.3.2	Diffrazione da fenditure multiple	100	14.4	La misura e il Principio d'indeterminazione	158
9.4	Qualche esempio pratico	101	14.5	Onde di materia	160
9.4.1	Diffrazione di onde sonore	101	14.6	Gli atomi	160
9.4.2	Diffrazione di onde luminose	102	14.6.1	Gli spettri atomici	161
9.4.3	Potere risolutivo	103	14.7	Quantizzazione del momento angolare	162
9.4.4	Diffrazione di raggi X	104	14.8	Lo spin degli elettroni	163
10	L'entropia	107	14.9	Il Principio di Pauli	165
10.1	Macchine termiche	107	14.9.1	La chimica	166
10.2	La Macchina di Carnot	109	14.9.2	Semiconduttori	169
10.3	Entropia	112	14.9.3	Il diodo	170
10.4	Il secondo principio della termodinamica	112	14.9.4	Il transistor	171
10.4.1	Passaggi di calore a volume costante	114	14.10	L'equazione di Schrödinger	172
10.5	L'espansione irreversibile di un gas	115	15	Una storia esemplare	177
10.6	Interpretazione microscopica dell'entropia	118	15.1	La scarica degli elettroscopi	177
10.7	La media e la varianza di una distribuzione	122	15.2	La scoperta dei raggi cosmici	179
11	La teoria della Relatività Ristretta	125	15.3	Caratteristiche dei raggi cosmici	179
11.1	Le trasformazioni di Lorentz	125	16	Chi l'ha ordinato?	183
11.2	La dilatazione del tempo	127	16.1	Particelle penetranti	183
11.3	Contrazione della lunghezza	129	16.2	L'ipotesi del neutrino	184
11.4	Composizione delle velocità	130	16.3	L'antimateria	186
11.5	I quadrivettori	131	16.4	La scoperta del muone	187
11.6	Il quadrivettore energia-impulso	132	16.5	La scoperta del pione	187
12	Muoversi tra sistemi	139	16.6	La lambda e i mesoni K	188
12.1	Una tecnica alternativa	140	17	I nuovi numeri quantici	191
12.2	Acceleratori e collider	143	17.1	I leptoni	191
13	La Relatività Generale	145	17.2	I barioni	191
13.1	La misura nei vari sistemi di riferimento	145	17.3	I mesoni	192
13.2	Il principio di equivalenza	146	17.4	Gli adroni	192
13.3	la geometria dell'Universo	148	17.5	Classificazione in base allo spin	193
13.4	Effetti gravitazionali sul tempo	150	18	Imitare la Natura	195
14	La Meccanica Quantistica	153	18.1	Gli acceleratori di particelle	195
14.1	Il corpo nero	153	19	Studiare le particelle	197
14.2	L'effetto fotoelettrico	155	19.1	Sezione d'urto	197
14.3	L'effetto Compton	157	19.2	Vita media	198
			20	Le risonanze	203
			20.1	Urti tra particelle	203
			20.2	La massa invariante	205

21 Le particelle strane	209	24.5 L'antimateria	230
21.1 I decadimenti della Λ	209	24.6 La produzione delle particelle strane	231
21.2 Produzione associata	211	24.7 L'interazione debole	231
22 Il Modello a Quark	213	25 Il bosone di Higgs	233
22.1 Tre nuove Tavole Periodiche	213	25.1 Richiami sul concetto di energia . . .	233
22.2 L'ipotesi dei quark	215	25.2 Campi autointeragenti	235
22.3 L'ottetto di mesoni	216	25.3 Sul significato dell'energia	235
22.4 L'ottetto di barioni	216	25.4 L'introduzione della relatività	236
22.5 Quark colorati	217	25.5 Il Meccanismo di Higgs	236
23 Il Modello Standard	221	25.6 Sulla forma del potenziale di Higgs .	237
23.1 I costituenti della materia	221	25.7 Campi massivi	239
24 Campi e Particelle	223	25.8 La massa dei bosoni vettori	240
24.1 Le forze fondamentali	223	Appendice	243
24.2 Una rivisitazione del concetto di energia	224	Approssimazione di funzioni	243
24.3 L'energia delle interazioni tra parti- celle	226	Equazioni differenziali a variabili separabili	245
24.4 Altri processi	229	Unità naturali	246
		Soluzione degli esercizi	249

Prefazione

Tra le parti di un libro la prefazione è sempre la meno letta. In questo caso è molto utile leggerla, sia per capire lo spirito con il quale è stato organizzato il materiale all'interno del volume, sia per comprendere le motivazioni alla base della scelta di redigere un testo così.

Il testo nasce dall'esperienza delle lezioni per gli studenti dei licei sulla fisica delle particelle, da me tenute nell'ambito dei programmi per l'orientamento e del Piano Lauree Scientifiche. Molti insegnanti, al termine delle mie lezioni, mi hanno chiesto materiale da utilizzare per riproporre in classe alcuni degli argomenti trattati, lamentando l'indisponibilità di testi adeguati. Per questo ho pensato di cominciare a scrivere queste note, con l'intento di ampliarle il più possibile nel corso del tempo, includendovi anche materiale più tradizionale.

Un'altra economia è possibile

Questo lavoro è una sfida al sistema economico corrente. La sfida si basa sul paradigma *Open Source*, da cui derivano le licenze *Creative Commons*.

Nell'attuale sistema economico il lavoro per la redazione di un libro, di gran lunga quello più faticoso e ricco di contenuto, è remunerato molto meno del lavoro necessario per la sua composizione, stampa, distribuzione e vendita. L'autore di un libro percepisce, mediamente, non più del 10–15 % del prezzo di copertina (a meno che non si tratti di un *best-seller*, naturalmente). Altrettanto percepisce il libraio, che fornisce il servizio di vendita. La maggior parte del prezzo di copertina va in compensi per l'editore e il distributore.

Grazie soltanto alla posizione di vantaggio, determinata dal posizionamento sul mercato e dalla rete di vendita, gli editori (che sono quelli che fanno

il prezzo) godono di un diritto economico non proporzionato al lavoro svolto. Questo è tanto più vero per i libri scolastici, per i quali il rischio d'impresa è assai basso, giacché l'editore conosce in anticipo l'ordine di grandezza del volume di vendite. In sostanza, pagando un libro di testo non si remunera il lavoro dell'autore, ma la posizione di vantaggio economico dell'editore e del distributore.

Questo stesso modello è adottato in moltissime attività economiche, dove ciò che determina il costo di una prestazione spesso non sono la quantità e la qualità del lavoro svolto, ma il possesso o meno di un presunto diritto a limitare la libertà dei clienti.

Noi crediamo, al contrario, che un altro modello economico sia possibile nel quale quel che deve essere remunerato è il lavoro i cui frutti, una volta remunerati, possano andare a beneficio dell'intera comunità. Che un tale modello sia concretamente attuabile è ampiamente dimostrato dal successo del software Open Source, che si può copiare, modificare, distribuire gratuitamente senza dover pagare *royalties* a nessuno. Le aziende che producono questo tipo di software, che dunque pagano gli stipendi dei programmatori e dei progettisti, funzionano e sono spesso più floride di quelle che vendono software tradizionale. Il business si basa sull'offerta del servizio, non sul possesso di un diritto a limitare la libertà altrui.

Siamo ormai così assuefatti al sistema che non ci rendiamo conto delle sue assurdità e siamo pronti a giustificarle con argomenti solo apparentemente ragionevoli. È il caso dei brevetti, ad esempio. Non è vero, come vuol farsi credere, che i brevetti aiutino lo sviluppo economico. È vero il contrario. Quando ci viene detto che l'industria farmaceutica, ad esempio, ha bisogno dei brevetti per coprire gli ingenti costi della ricerca, si tratta palesemente di una bufala.

I costi della ricerca, infatti, sono coperti dagli Stati e dai cittadini attraverso l'acquisto dei medicinali. Il costo della ricerca è alto perché ciascuna industria deve svolgere, segretamente, le stesse attività delle altre, con una moltiplicazione degli sforzi enorme e un costo esorbitante. Mettendo in comune i risultati, ogni azienda potrebbe risparmiare miliardi di euro e i farmaci potrebbero costare molto meno. Se fossero gli Stati a coprire questi costi, e i risultati fossero pubblici, il prezzo dei farmaci scenderebbe in maniera consistente e alla fine ci sarebbe un risparmio per tutti. Lo stesso vale per altri settori. È del tutto evidente, ad esempio, che i brevetti non servono alle industrie per acquisire una posizione di vantaggio rispetto ai potenziali concorrenti. Se così fosse i prodotti innovativi dovrebbero essere appannaggio di una sola azienda, ma non è così (pensate solo all'industria dei tablet: un'azienda ne ha lanciato uno, coperto da brevetto, e tutte le altre l'hanno seguita a ruota, aggirando il brevetto o acquistandolo). I brevetti hanno il solo scopo di creare un mercato delle idee innovative che, se ci si pensa bene, è un mercato del tutto irragionevole e contrario a ogni principio etico.

Noi pensiamo che l'autore di un libro di testo vada remunerato per il suo lavoro e non per l'aver acquisito una certa posizione di mercato. Una volta redatto il libro e una volta che il compenso per l'autore sia stato equo, il libro deve poter essere distribuito quasi gratuitamente: si dovrebbero pagare solo i costi effettivi della sua distribuzione e il giusto compenso per coloro che la rendono possibile. Non c'è ragione per cui un libro stampato da anni, che abbia già venduto migliaia di copie, non possa essere fotocopiato o reso pubblico. Sia ben chiaro che il modello che proponiamo non chiede di rinunciare alla proprietà intellettuale: il **diritto d'autore** resta di esclusiva proprietà dello stesso ed è inalienabile. È il **diritto esclusivo di sfruttamento economico** delle opere che, nella legislazione corrente, è cedibile ad altri e che noi riteniamo sia quanto meno da modificare. Per questo la licenza d'uso di questo libro impone che si citi sempre l'autore ogni volta che se ne fa un uso pubblico.

Questo modello non impedisce l'esistenza di case

editrici, che possono (devono) basare il loro business sull'efficacia della distribuzione, sul valore aggiunto, sulla capacità di offrire servizi diversi. La nascita di **WIKIPEDIA** non impedisce agli editori di vendere enciclopedie e dizionari. Ne modifica, evidentemente, il profilo. Chi usa le informazioni reperite online per acquisire un'informazione non produce un danno all'editore più di quanto non faccia una signora che acquisti da una bancarella una borsa firmata da un grande stilista nei confronti di quest'ultimo. La signora non avrebbe comunque mai acquistato quella borsa al prezzo proposto dallo stilista.

Si potrebbero inventare decine di modelli economici alternativi basati su un paradigma aperto, ma questo dovrebbe essere lavoro per economisti. Noi qui lanciamo la sfida. Rendiamo pubblico questo testo iniziale, chiedendo un supporto economico volontario a coloro che decideranno di adottarlo. Se riusciremo ad accumulare una cifra ritenuta ragionevole, quale compenso per questo lavoro, ne realizzeremo un altro (per altri gradi dell'istruzione scolastica o introducendo nuovi argomenti e nuove tecnologie). Se avremo successo e i proventi saranno sufficienti, potremo remunerare il lavoro di altri professionisti per la realizzazione di filmati, animazioni e altri supporti multimediali, che in questo caso sono stati tutti realizzati dall'autore, in prima persona, con evidente dispendio di energie.

Se condividi questa visione del mondo e ti sembra che il contenuto del libro sia adatto alle tue esigenze (e quest'ultimo è il requisito più importante), diffondilo e invita a supportarlo. Ti invitiamo anche a inviarci commenti, segnalazioni su possibili errori, suggerimenti, sia sul contenuto, sia sul modo di presentarli, sia sul modello distributivo. Naturalmente non garantiamo l'accoglimento di tutti i suggerimenti che potranno pervenire, perché di nuovo questo fa parte della libertà di ciascuno di realizzare le proprie opere come crede. Il che, però, non impedisce agli altri, una volta venuti in possesso di tali opere, di fare altrettanto.

Il titolo

Il titolo di questo volume non è stato scelto a caso. L'Italiano è una lingua che si presta a diverse, interessanti, e talvolta divertenti, interpretazioni del significato delle parole. In particolare l'aggettivo **sperimentale** utilizzato nel titolo ha in questo testo significati diversi, tutti contemporaneamente validi.

È sperimentale, come abbiamo detto sopra, il modo in cui il testo è realizzato e distribuito. Si tratta, cioè, della sperimentazione, della ricerca di un nuovo modello economico.

L'aggettivo sperimentale si riferisce anche al taglio dato all'introduzione dei concetti della fisica. Molti testi di fisica appaiono più come libri di matematica, nei quali si danno certe definizioni allo studente e se ne traggono le conseguenze. Le definizioni, in molti casi, piovono dall'alto, senza una spiegazione plausibile sul perché sia il caso di introdurle o su quale sia la loro ragion d'essere. In questo testo la fisica viene introdotta attraverso l'esperimento. Ogni argomento viene analizzato a partire dalle osservazioni sperimentali, che determinano le grandezze fisiche d'interesse, portando naturalmente alla formulazione delle leggi fisiche.

Infine, è sperimentale il mezzo scelto per la realizzazione del testo. Il supporto elettronico consente di fruire di contenuti multimediali e delle potenzialità dell'ipertesto. Si potrebbe fare molto di più, in effetti. La tecnologia è matura. Ma, spesso a causa di scelte determinate dal modello economico di cui parliamo sopra, molti produttori di software non consentono di usare in maniera semplice le innovazioni disponibili. Naturalmente il problema si potrebbe superare realizzando *ad hoc* anche i lettori per il supporto, ma questo avrebbe un costo eccessivo per noi (almeno in questa fase) e in ogni caso limiterebbe la platea di potenziali fruitori dell'opera. Possiamo solo sperare che il sistema avrà successo e ci consentirà, in futuro, di aumentare sempre di più l'offerta.

Tecnologia

Per realizzare questo prodotto sono stati usati per lo più strumenti *Open Source* o con licenza *Creative Commons*.

Il testo è stato redatto con $\text{\LaTeX}$ ¹: un programma per la composizione di testi estremamente potente e liberamente scaricabile dalla rete, basato sul suo predecessore $\text{\TeX}$ sviluppato da Donald Knuth, Professore di *Computer Science* a Stanford e autore di una monumentale opera sulla programmazione dei computer [1].

I filmati sono stati editati con *MPEG Streamclip* e *iMovie*. Per creare o manipolare alcune figure è stato usato *Gimp*.

Le musiche sono state scaricate da *jamendo*: un sito che raccoglie musica con licenza *Creative Commons*.

Le animazioni sono state realizzate con *Animation Desk* per iPad.

L'immagine di copertina è di Alegri (*alegriphotos.com*).

Alcuni *link* come [questo](#) potrebbero non portare a nulla (né a una pagina web, né a un altro punto del testo). Questo perché il link è previsto portare il lettore a una sezione del testo in cui si parla dell'argomento relativo che non è ancora stato prodotto. Molti link porteranno a pagine di *WIKIPEDIA* (in inglese). Abbiamo scelto di usare questo strumento quale riferimento a informazioni non strettamente pertinenti l'argomento del testo, nonostante le critiche che vengono da più parti sull'attendibilità delle informazioni che vi si trovano. L'opinione dell'autore è che questo strumento, come tutti, evolverà col tempo diventando sempre più attendibile. Riguardo le critiche relative al fatto che le informazioni presenti sul sito sono copiate da altre fonti senza verifica, va detto che lo stesso è accaduto e accade tuttora con i libri e che solo in rari casi gli autori dei libri verificano le informazioni sulle fonti originali (noi, naturalmente, l'abbiamo fatto ove possibile).

La non apertura dei sistemi operativi e delle applicazioni continua purtroppo a provocare inconve-

¹ $\text{\LaTeX}$ si pronuncia *latek*

nienti piuttosto deprecabili. Uno di questi consiste nel fatto che la versione elettronica di questo testo si può usare senza alcun problema su un computer, mentre sui tablet esistono limitazioni (incomprensibili). Ad esempio, lo stesso **Acrobat Reader** su iPad non supporta la visione dei filmati *embedded*. Lo stesso accade con **iBooks**. Secondo la nostra esperienza i filmati sono invece perfettamente visibili se si carica il PDF su **Dropbox** e si visualizza usando il browser interno dell'applicazione per tablet (tuttavia in questo caso non si può avere il testo in modalità *full screen*), oppure usando **ezPDF Reader**². Confidiamo che il mercato spinga nella direzione giusta e che sia possibile usare un'applicazione in modo uniforme su ogni piattaforma.

Formare, non informare

Il semplice trasferimento di conoscenza non ha molto senso. Conoscere le leggi della fisica è utile, ma non indispensabile nella vita di una persona, tanto meno se questa conoscenza si limita alla mera capacità di scrivere le formule corrispondenti senza capirle.

Capire le leggi della fisica e il processo che ha condotto alla loro formulazione, al contrario, è di fondamentale importanza per la formazione complessiva degli studenti. Ecco perché questo testo pone l'accento più sul **come** si arrivi a formulare le leggi fisiche piuttosto che su queste ultime. In particolare, le leggi fisiche davvero fondamentali sono poche ed è su queste che si concentra tutta la struttura del volume. Le leggi derivate da quelle fondamentali sono trattate come esercizi e non come parte integrante del testo. Questo non vuol dire che si possano ignorare, ma che non si devono necessariamente *ricordare*. Laddove esistano relazioni particolari che vale la pena siano ricordate a memoria per la frequenza con la quale si usano o per l'importanza che rivestono nel loro ambito, queste sono evidenziate in **rosso**, anche negli esercizi.

La matematica presente in ogni parte del volume (a parte gli esercizi) è ridotta al minimo indi-

spensabile e non si assume la conoscenza di concetti avanzati, in modo tale che il testo possa essere usato da scuole diverse (Licei scientifici, classici, scuole professionali).

Come usare questo testo

Il formato migliore per questo testo è quello elettronico (è per questo che è nato). La versione **portrait** può essere più comoda per i laptop e per alcuni tablet. Puoi comunque stampare il testo (anche se in questo caso perderai la funzionalità dei filmati, che tuttavia trovi sempre sul **canale YouTube** dell'autore). La versione **portrait** è più adatta per la stampa. Puoi fare la stampa da te o rivolgerti a un servizio specializzato (una copisteria, una tipografia o su Internet).

Se pensi di usare questo testo a scuola in formato cartaceo, puoi scaricarne una copia che puoi fornire a un tipografo di zona perché ne stampi un numero sufficiente di copie. Gli alunni possono quindi acquistare le copie versando solamente il corrispettivo relativo al costo vivo della stampa, che in questo caso probabilmente è inferiore alla spesa necessaria per stampare il testo in proprio. La licenza con cui è distribuito questo testo consente a chiunque (anche a un istituto scolastico) di stampare le copie necessarie e di fornirle agli studenti a un prezzo superiore a quello corrispondente al costo dell'operazione. In questo modo la Scuola può percepire un contributo extra utile in questo periodo di tagli. Va da sé che se il ricavo è eccessivo, gli studenti preferiranno procurarsi il materiale da soli: il modello di distribuzione che abbiamo scelto consente a più persone di creare valore, ma senza eccessi. Qualunque abuso sarebbe automaticamente eliminato dal *mercato* dalla disponibilità gratuita del bene.

Ovviamente sei sempre invitato a versare un contributo per l'autore. In questo modo avrai la garanzia che il prodotto sarà sempre mantenuto al meglio e continuamente migliorato. Potrai anche sostenere il modello di sviluppo scelto, dimostrando che un mercato equo e sostenibile è possibile, nel quale è il lavoro o un servizio a essere retribuito e non una

²A pagamento, ma dal prezzo accessibile.

rendita di posizione o da capitale. Nel caso in cui il testo sia adottato a scuola, il metodo migliore per versare il contributo consiste nel raccogliere il denaro che avete deciso di versare e fare un unico versamento. Non è necessario che il contributo sia lo stesso per tutti gli studenti. Poiché il versamento è libero (si tratta di una donazione) puoi prevedere la possibilità che alcuni studenti (quelli le cui famiglie hanno problemi di carattere economico) non paghino per l'uso di questa risorsa.

Per l'insegnante

Il testo contiene molto più materiale rispetto a quello che si può normalmente pensare d'insegnare alla maggior parte degli studenti. La lunghezza del testo non deve spaventare: abbiamo scelto di spendere molte parole perché crediamo che della fisica si debba insegnare sopra tutto il metodo e non tanto i contenuti, pure indispensabili. È importante capire il significato delle equazioni e il modo in cui si ricavano. Non è quindi il numero di pagine che suggerisce di limitare gli argomenti, ma il fatto che oggettivamente alcuni sono molto (troppo) difficili per molti studenti. Lasciarli però consente agli insegnanti di preparare gli argomenti da trattare in maniera più consapevole e completa, e agli studenti più bravi di approfondire da soli argomenti che altrimenti sarebbero rimasti troppo vaghi.

In alcuni capitoli la matematica può essere pesante o inadeguata rispetto alla preparazione dello studente al momento in cui l'argomento è proposto: questo non è un problema. In tutti i casi l'argomento si può discutere senza entrare troppo nei dettagli formali e se ne può proporre una lettura più qualitativa, aiutandosi con figure, esperimenti e simulazioni (ad esempio, nei capitoli riguardanti le onde è richiesta la conoscenza della trigonometria, ma gli stessi argomenti si possono affrontare impiegando unicamente le simulazioni proposte e facendo qualche esperimento).

I capitoli possono contenere indicazioni sui prerequisiti. In generale tenderemo a ridurre al minimo i prerequisiti in modo che l'insegnante possa seguire un suo proprio percorso, senza essere costretto a

seguire un ordine preciso. In questo testo per prerequisiti si intende la conoscenza di argomenti necessaria per comprendere il formalismo impiegato. La descrizione qualitativa di certi fenomeni non richiede conoscenze pregresse, perché si può sempre dare per nota. Ad esempio, nel caso del paragrafo sul **corpo nero**, il fatto che le particelle cariche accelerate irraggiano onde elettromagnetiche non è considerato un prerequisito, anche se il fenomeno viene citato. Se il fenomeno non è stato trattato, si può fornire questa informazione agli studenti che possono acquisirla come vera anche se non la comprendono a fondo. In questo caso basta *fidarsi* dell'insegnante, rimandando la comprensione del fenomeno a un secondo momento.

Piano dell'opera

L'autore s'impegna, a fronte di un numero sufficiente di donazioni, ad ampliare i contenuti al maggior numero possibile di campi della fisica (inclusa la fisica classica che sarà trattata sempre con un approccio di tipo sperimentale). La sequenza con la quale nuovi moduli saranno resi disponibili non sarà necessariamente coerente con quella nella quale gli argomenti vengono usualmente presentati, né con quella che l'autore ritiene la migliore sequenza possibile. La sequenza sarà dettata per lo più dalle richieste di studenti e insegnanti che avranno apprezzato i moduli precedenti.

Alla scoperta dell'Universo

Capire come funziona l'Universo è il compito primario di un fisico. Chi non intraprende questa carriera può senz'altro nutrire legittima curiosità circa questo argomento, ed è per questa ragione che studierà un po' di fisica. La curiosità, in fondo, è sempre stata il propulsore della scienza, in ogni suo aspetto. La ricerca scientifica, in effetti, si fa solo per questa ragione, non certo perché gli scienziati pensano di poter offrire ai loro concittadini qualche nuovo prodotto tecnologico inteso a migliorare la loro vita o semplicemente a renderla più divertente! Non si chiede mai a un musicista o a uno scrittore qual è l'utilità pratica della sua opera, ma chissà per quale motivo a chi fa ricerca scientifica si rivolge spesso questa domanda. Chi fa questo mestiere lo fa con lo stesso spirito con il quale un musicista compone o uno scrittore scrive. La spinta a farlo non ha alcuna finalità pratica e nasce dall'esigenza di comunicare qualcosa agli altri.

Non accade spesso, ma le opere di alcuni musicisti e scrittori possono produrre effetti molto concreti: pensate soltanto al giro d'affari conseguente alla vendita dei dischi e all'organizzazione di concerti, che non apporta benefici economici solo al compositore, ma anche al produttore, agli operatori dell'industria discografica, venditori, informatici, trasportatori, operai e decine di altre categorie (e poi c'era qualcuno che diceva che con la cultura non si mangia!). Una volta giunti nelle case dei fruitori i prodotti possono dar loro piacere, si possono impiegare per animare feste o riunioni, per sottolineare momenti speciali, etc.. Se non ci fossero stati i musicisti, gli scrittori, i pittori, gli scultori e tutti gli altri produttori di *nulla*, l'umanità non sarebbe certo al livello cui si trova oggi!

Lo stesso vale per la fisica e tutte le altre scienze. Il motivo per il quale si deve studiare la fisica, in fon-

do, è lo stesso per cui si studiano Dante e Petrarca, Verdi e Beethoven, Raffaello e Picasso.

Anche nel caso della scienza accade che, sebbene l'obiettivo primario sia quello di capire il funzionamento dell'Universo, talvolta l'acquisizione di queste conoscenze porti alla realizzazione di qualcosa di utile. Abbiamo il frigorifero nelle nostre case perché qualcuno ha studiato, in passato, la natura del calore; i telefoni cellulari grazie all'elettromagnetismo e i navigatori GPS devono la loro capacità di guidarci nel punto in cui siamo diretti alla relatività di Einstein; tutti i dispositivi elettronici funzionano grazie alla meccanica quantistica e potremmo andare avanti per molte pagine.

Anche chi non ha particolari curiosità dunque vorrà studiare un po' di fisica, se non altro per capire un minimo il mondo che lo circonda e, in qualche caso, fare qualche scelta in modo più consapevole.

Quel che è più importante per chi studia fisica (senza avere l'obiettivo di diventare un fisico) non è, naturalmente, apprendere i contenuti di questa disciplina, allo stesso modo di come lo studio della Divina Commedia non comporta per ciascuno la conoscenza dettagliata (magari a memoria) del contenuto di tutti i Canti dell'opera. Prima di conoscere nel dettaglio tutta l'Opera bisogna conoscerne sommariamente i contenuti e la genesi, non trascurando qualche particolare sull'Autore. Conoscere le Leggi della gravitazione universale, i principi dell'elettromagnetismo, la termodinamica, la meccanica quantistica e la relatività ristretta e generale non è altrettanto importante quanto conoscere il **metodo**. Il metodo della fisica è il **metodo sperimentale**, che consiste nel fare osservazioni e misure e, da queste, ricavare modelli matematici che descrivano i fenomeni osservati, ma che siano anche capaci di permettere al fisico di fare **predizioni** su quanto

potrà osservare. È importante osservare che in questo contesto una predizione non è necessariamente qualcosa che si deve poter fare in anticipo rispetto a un'osservazione. Una predizione è, in fisica, il risultato di un esperimento che si può ottenere attraverso la **teoria**. L'esperimento può anche essere stato fatto in precedenza. Una **teoria** è un insieme logicamente coerente di osservazioni sperimentali da cui si può derivare un modello con caratteristiche predittive: un modello che non prenda le mosse da un fatto sperimentale non si può definire una teoria. Allo stesso modo un modello che non faccia previsioni non si può definire una teoria. Infine, un modello che non porti a predizioni differenti rispetto ad altri non è una teoria diversa dalla precedente: è solo un modo alternativo di formulare una stessa teoria. Qualche esempio chiarirà il significato di quanto esposto.

Predizioni

La capacità di una teoria di fare una predizione non si misura dalla sua capacità di prevedere cosa succederà in un futuro più o meno lontano a un certo sistema. È ben noto, ad esempio, che non si possono prevedere i terremoti, ma la teoria della tettonica a zolle permette di *predire* i luoghi in cui questi avverranno con maggiore frequenza, avendo elaborato un modello nel quale le placche della superficie terrestre sono in continuo movimento e per questa ragione possono andare incontro a fenomeni particolarmente violenti dovuti alle tensioni accumulate nelle fasi di scontro.

Quando Einstein formulò la sua teoria della Relatività Generale trovò che un pianeta come Mercurio avrebbe dovuto percorrere un'orbita il cui perielio ruotava attorno al Sole più rapidamente di quanto non dovesse fare per effetto delle forze gravitazionali. In effetti all'epoca si sapeva già da tempo che il perielio di Mercurio precedeva attorno al Sole più rapidamente di quanto predetto dalla teoria della gravitazione di Newton, ma non era stata trovata, fino ad allora, alcuna spiegazione convincente. La teoria di Einstein, al contrario, *prevedeva*

una velocità di precessione identica a quella misurata sperimentalmente ben prima che la teoria fosse formulata.

Teorie e fatti sperimentali

Nessuna tra le teorie della fisica può fare a meno di fatti sperimentali. Questi sono indispensabili per poterle formulare. Una *teoria* che ipotizzi l'esistenza di qualche relazione tra grandezze fisiche in virtù di pretese evidenze di tipo logico, non è una teoria che un fisico può prendere in considerazione. Per esempio, non si può sostenere che l'Universo deve essere il risultato della volontà di un Creatore, perché non esistono fatti sperimentali che supportino una tale affermazione. Né quest'affermazione si può configurare come una conseguenza di un fatto sperimentale. Naturalmente, questo non vuol dire che **sicuramente** l'Universo non è stato fatto da un Creatore: vuole soltanto dire che non è la fisica a occuparsi di questo problema. Un fisico può credere o meno all'esistenza di un Creatore, senza che questo turbi qualcuno.

Sebbene la questione sia molto controversa, secondo il parere dell'autore, l'**omeopatia** non si può considerare una disciplina scientifica perché prende le mosse da un assunto, un principio, formulato da un medico del XVIII secolo³ senza alcuna base sperimentale, ma solo come un'ipotesi possibile. L'ipotesi che per curare le malattie si possano usare sostanze che nelle persone sane inducono sintomi simili a quelli osservati nelle persone malate forse può essere considerata da qualcuno plausibile, ma di certo non è qualcosa che si può osservare come un fatto sperimentale a partire dal quale si può enunciare una teoria!

Viceversa, la teoria della relatività speciale, i cui risultati appaiono ai più privi di senso (ma a pensarci bene vale esattamente il contrario: sono i risultati della meccanica classica a essere privi di senso in certe condizioni), è una teoria perché prende le mosse da un fatto sperimentale tanto semplice quanto

³Samuel Hahnemann, vissuto tra il 1755 e il 1843.

apparentemente privo di logica: la velocità della luce è la stessa qualunque sia lo stato di moto relativo dell'osservatore!

Teorie e previsioni

Una teoria che non faccia previsioni non è una teoria: è solo un'ipotesi sulla quale si può lavorare cercando di trovare il modo di estrarne delle previsioni, ma fino a quando questo non è possibile non si può parlare di teoria.

Uno dei risultati sperimentali della fisica moderna è che la materia è composta di un numero relativamente grande di particelle elementari: sei *leptoni* (tra cui l'elettrone), diciotto *quark* e tutte le corrispondenti antiparticelle. Una struttura così complessa fa indubbiamente pensare alla possibilità che queste 48 particelle non siano affatto elementari, ma che siano a loro volta composte di altre particelle. In assenza di previsioni su cosa si dovrebbe osservare in un esperimento che tenti di confermare o meno quest'ipotesi, l'ipotesi stessa non si può definire teoria.

Teorie alternative

Due teorie che facciano entrambe le stesse previsioni non sono due teorie: sono la stessa teoria formulata in modo diverso. I risultati della meccanica quantistica, per esempio, si possono ottenere sia con la meccanica ondulatoria che con quella delle matrici: due formalismi matematici molto diversi, che hanno in comune gli stessi dati sperimentali e fanno le stesse previsioni. Non si può dire quale sia la teoria corretta: lo sono entrambe nel limite in cui riescono a fare predizioni verificabili sperimentalmente.

Capire la Fisica

Capire la Fisica significa capire quali osservazioni e quali ragionamenti hanno portato alla formulazione dei modelli che usiamo oggi per comprendere quel

che accade nell'Universo. A ben vedere non è questo quel che fa la maggior parte dei libri di testo per le scuole superiori e nemmeno per l'Università. L'approccio seguito in questo testo è completamente differente da quello della maggior parte dei libri in circolazione (senza che questo implichi un giudizio di merito: non è detto che l'approccio di questa pubblicazione sia il più efficace; è solo un approccio alternativo). Non pretende d'insegnare la Fisica per **temi**, ma cerca di portare il lettore a capire come siamo giunti alle nostre attuali conoscenze illustrando un **metodo sperimentale**. Non lo fa attraverso una narrazione storicamente attendibile dei fatti che hanno portato all'attuale formulazione della Fisica, ma in un certo senso *raccontando* un percorso immaginario che avrebbe potuto portare agli stessi risultati, come fosse un romanzo che attraverso la finzione pretende d'illustrare la realtà. Anche per questo il volume reca il titolo **Fisica Sperimentale**: perché la sequenza scelta per introdurre i diversi temi è inusuale e funzionale a questa *narrazione*.

Non si tratta di una semplificazione: tutt'altro! L'operazione è complessa e richiede un certo sforzo per poter essere compresa. È più facile affidarsi all'autorità di qualcuno e *credere* che quel che dice è vero. Con questo libro vogliamo convincervi che l'attuale visione del mondo che hanno i fisici non è il risultato dell'adesione a un modello confezionato da un'Autorità cui tutti dobbiamo obbedienza, ma una logica conseguenza dei fatti sperimentali. Il bello della scienza è che non esiste alcuna Autorità: Albert Einstein non credeva alla Meccanica Quantistica, né al fatto che l'Universo sia in espansione, ma per quanto ne sappiamo oggi aveva torto. Sarà anche stato uno degli scienziati più influenti e profondi del secolo scorso, ma questo non gli dava e non gli dà il diritto di decidere cosa sia giusto e cosa no: questa decisione dipende esclusivamente dai **fatti sperimentali**. Il percorso dunque non è facile e, per dar modo a tutti di poter apprezzare fino in fondo tutte le conseguenze di ciascuna scoperta, il volume è corredato di dettagliate analisi dei fenomeni, anche di tipo formale e matematico. Per capire la fisica non è però necessario affrontare tutti gli aspetti di ogni tema, né tantomeno approfondirli

fino al grado di dettaglio proposto (che comunque è ridotto). Se sei uno studente, l'insegnante potrà guidarti nella selezione degli argomenti che riterrà più opportuna. Se sei un insegnante potrai trovare qualche spunto per presentare in modo diverso certi temi. Anche la sequenza degli argomenti si potrà modificare secondo le esigenze. La struttura dei capitoli è pensata per essere alterata quasi a piacimento: ogni capitolo è scritto in modo tale che sia il più possibile indipendente dagli altri. Anche se in certi casi si richiamano argomenti trattati in altri capitoli, in ciascuno di essi ci si può affidare, in questi casi, all'Autorità e *credere* (temporaneamente) a quanto affermato. Non bisogna mai dimenticare, però, che tutto quel che è scritto in questo volume deve essere sempre guardato con il massimo sospetto e non bisogna mai credere fino in fondo a quel che è stato scritto finora su libri e articoli di qualsiasi disciplina scientifica. Solo così si può avere la speranza di contribuire in qualche modo al progresso scientifico. Se chi ci ha preceduto avesse *creduto* fermamente nella correttezza degli argomenti avanzati dai propri maestri, oggi saremmo lontanissimi dal sapere come funziona l'Universo.

Scoprire la Fisica

Questo libro non vi condurrà a imparare le Leggi della Fisica, ma a **scoprirle**, partendo dai fatti sperimentali attraverso la formulazione di ipotesi ragionevoli che leghino questi fatti a un modello plausibile della realtà che per un fisico non è nient'altro che ciò che si misura. Non ha molto senso chiedersi se esista una realtà diversa da quella da noi percepita (per lo meno non lo ha per un fisico; forse per un filosofo un senso ce l'ha).

In questo capitolo vi forniamo un primo, molto sommario riassunto del percorso che intendiamo proporvi per affrontare lo studio della fisica. Si tratta di un percorso indubbiamente originale, ma che vi farà vedere i diversi aspetti della fisica come fortemente collegati, quali in effetti sono.

1.1 Il problema della misura

Per iniziare a parlare di Fisica dobbiamo prima di tutto imparare a **misurare** qualche **grandezza fisica**. Tra le diverse misure che possiamo fare una è particolarmente interessante perché non banale: la **temperatura**. Lo studio delle variazioni di temperatura ci conduce a ipotizzare l'esistenza di *qualcosa* che transita da un corpo all'altro cui diamo il nome di **calore** la cui natura tuttavia appare misteriosa. Per capire cos'è il calore possiamo osservare che la temperatura di un corpo cambia anche se non viene posto in contatto con un altro corpo a temperatura diversa: è sufficiente esporlo alla luce per riscaldarlo.

1.2 La luce

Se vogliamo capire cos'è il calore dobbiamo dunque studiare le proprietà della luce e per questo affrontiamo lo studio dell'ottica iniziando dall'**ottica geometrica**.

Dall'analisi delle proprietà dei raggi luminosi si può concludere che la luce potrebbe essere composta di un flusso di particelle emesse continuamente dalle sorgenti luminose. Questi corpuscoli, urtando certe superfici, possono *rimbalzare* e dar luogo al fenomeno della riflessione, e possono cambiare direzione passando da un tipo di mezzo a un altro dando luogo così al fenomeno della rifrazione. D'altra parte, se la luce è fatta di corpuscoli che viaggiano in direzione dei raggi luminosi, si può spiegare un'altra osservazione sperimentale: i corpi si scaldano anche se sono colpiti da altri corpi in movimento (pensate a un chiodo la cui testa è picchiata con un martello: sia il martello che il chiodo si scaldano). È facile pensare, così, che il calore prodotto dalla luce sia il risultato del *picchiare* dei corpuscoli di luce sui corpi illuminati.

1.3 Le onde

A pensarci bene, però, c'è un'altra spiegazione possibile: la luce potrebbe essere semplicemente il risultato di una perturbazione che si propaga come un'**onda** in un mezzo. In effetti le onde subiscono sia la riflessione che la rifrazione. Come si fa a decidere? È semplice: si escogita un esperimento che evidenzia la natura della luce cercando di trovare fenomeni che diano risultati diversi nel caso in cui la luce sia fatta di corpuscoli o di onde. Un tale espe-

rimento esiste e consiste nell'ostacolare il moto con uno schermo sul quale siano state praticate una o più **fenditure**: nel caso in cui ciò che si propaga al di qua dello schermo sia un flusso di corpuscoli, al di là di esso ci si aspetta di vedere gli stessi corpuscoli attraversare le fenditure; nel caso delle onde, invece, ci si aspetta di osservare il fenomeno della **diffrazione**, che consiste nella modifica della forma del fronte d'onda.

Gli esperimenti suggeriscono inequivocabilmente che la luce sia fatta di onde, quindi il fatto che riescano a fornire calore ai corpi illuminati non è spiegabile in termini di urti tra corpuscoli luminosi e corpi illuminati.

D'altra parte che il calore sia legato al moto è evidente: non solo un chiodo colpito da una martellata si scalda; lo fa anche qualunque oggetto sottoposto all'azione di qualche altro oggetto che si muove rispetto al primo. Per esempio, se si passa della carta vetrata su un pezzo di metallo si scaldano entrambi, ma anche se ci si frega le mani l'una contro l'altra o se si fa un buco sul muro usando un trapano.

La scoperta che la luce è un'onda non altera sostanzialmente quest'ipotesi. Il propagarsi di un'onda implica il moto del mezzo nel quale l'onda si propaga: sul mare il propagarsi delle onde provoca la variazione dell'altezza della superficie dell'acqua che va continuamente su e giù; nel caso della propagazione di un'onda sonora quel che si muove è l'aria presente tra la sorgente e chi ode il suono (togliendo questa non si ode più alcun suono). Di fatto, anche nel caso in cui assumessimo che la luce sia un'onda possiamo sempre pensare che gli effetti termici siano in qualche maniera il risultato del moto che le onde inducono sui corpi colpiti.

1.4 Il moto, il lavoro e l'energia

A questo punto non resta che studiare il **moto** e le sue cause, per capire cosa generi il calore. Lo studio del moto ci condurrà alla scoperta del concetto di **forza**. Per scaldare due oggetti in moto relativo è necessario che si manifesti una forza tra essi. Nel caso della carta vetrata sul blocco di metallo, questa

forza è l'attrito, così come nel caso della punta del trapano che penetra nel muro. Si può anche scaldare dell'acqua facendo muovere delle pale immerse in essa, azionate da qualche tipo di forza (la gravità, che fa cadere gli oggetti o quella di un motore a scoppio o elettrico). In ogni caso è evidente che la quantità di calore che si riesce a cedere agli oggetti in questo modo dipende dall'intensità della forza applicata, ma anche dallo spostamento relativo tra i due oggetti sfregati (se la carta vetrata non si muove rispetto al blocco non si produce alcun calore; il calore prodotto è tanto maggiore quanto più è lungo lo spostamento). Quest'osservazione c'indurrà a definire una nuova grandezza fisica che rappresenti questa caratteristica: il **lavoro**.

1.5 La termodinamica

Si vedrà presto che lavoro e calore sono due aspetti diversi della stessa grandezza fisica: l'**energia**. Trasferire calore o fare del lavoro significa semplicemente cambiare l'energia di un sistema, il che può avvenire attraverso una modifica delle velocità delle componenti del sistema (che determinano quella che si chiama **energia cinetica**) o per mezzo di una modifica delle relazioni spaziali tra particelle tra loro interagenti o sottoposte all'azione di forze esterne (**energia potenziale**). Il sistema più semplice che possiamo pensare di analizzare è evidentemente un sistema in cui queste forze siano nulle, per cui le particelle che lo compongono sono libere di muoversi nel volume nel quale sono contenute e questo porta a considerare lo studio dei **gas** come il sistema migliore per capire la natura del calore.

Dall'analisi delle proprietà di un gas si comprenderà come la temperatura sia semplicemente una misura della velocità media delle particelle che compongono i corpi (se un gas è fatto di particelle, devono esserlo anche i solidi e i liquidi, l'unica differenza essendo dovuta alle diverse forze con le quali i componenti interagiscono tra loro).

1.6 Campi di forze

Facendo gli esperimenti per il trasferimento di calore attraverso lo strfinio potremmo aver notato che, in certi casi, i corpi sottoposti a questo trattamento manifestano la capacità di attrarre altri oggetti o di attrarre o respingere oggetti della stessa natura anch'essi sfregati opportunamente contro un altro corpo. Le forze che si manifestano in questi casi non prevedono alcun tipo di contatto tra gli oggetti, un po' come avviene per quelle responsabili della caduta dei corpi. Giungiamo così all'idea di **campo** e dell'azione a distanza.

Dopo una panoramica sulle proprietà generali dei campi potremo così affrontare lo studio delle forze che agiscono a distanza, come le forze **gravitazionali** e quelle **elettriche**. Lo studio delle forze che sperimentiamo quotidianamente responsabili della caduta degli oggetti ci permetterà di capire la **gravità** e il moto dei corpi celesti. Lo studio delle forze di natura elettrica ci condurrà a ipotizzare che, se i corpi sono costituiti, come sembra, di corpuscoli molto minuti, almeno una parte di essi deve essere elettricamente carica. Attraverso lo studio dei fenomeni elettrici capiremo che tutti i fenomeni osservati si possono giustificare con quest'ipotesi e ne concluderemo che la materia dev'essere formata da particelle di carica elettrica positiva e negativa, oltre che di particelle prive di carica.

1.7 Correnti elettriche

Per capire come funzionano i campi elettrici utilizzeremo strumenti (i condensatori e le pile) che ci permetteranno di osservare un nuovo fenomeno: quello del passaggio della **corrente elettrica**, che si può sempre spiegare in termini di moto dei corpuscoli che costituiscono la materia dei conduttori.

1.8 Campi magnetici

Lo studio del comportamento delle correnti ci farà fare ancora un'altra scoperta: l'esistenza di una

forza nuova che si manifesta tra fili percorsi da corrente. Interpretando questo fenomeno alla luce della nostra ipotesi sulla costituzione della materia potremo verificare che in effetti le particelle elettricamente cariche in moto sono soggette a questa forza che definiamo magnetica alla quale si può associare un campo. Sarà quindi opportuno studiare anche questo campo.

1.9 Corpi rigidi

I fili conduttori percorsi da corrente sono soggetti a forze dovute alla presenza di campi magnetici e si muovono in seguito al manifestarsi di queste forze. Ma i fili non sono oggetti puntiformi: sono oggetti estesi per i quali ogni loro parte è legata alle altre da forze che fanno in modo che il filo mantenga la sua forma. Il problema di comprendere il moto di questi oggetti complicati appare altrettanto complicato, ma in realtà vedremo come sia possibile darne una descrizione generale tutto sommato molto semplice.

1.10 Onde elettromagnetiche

Le particelle elettricamente cariche hanno la proprietà di essere soggette a forze di natura elettrica e magnetica e, a loro volta, di essere sorgenti di tali forze. Lo studio del moto di queste particelle, indotto dalla presenza di **campi elettromagnetici** permetterà di stabilire che esse stesse devono essere sorgenti di campi elettromagnetici variabili nel tempo che, oltre ad avere carattere ondulatorio, possiedono tutte le altre caratteristiche della luce. È immediato, a questo punto, identificare la luce con il campo elettromagnetico.

1.11 Ancora sulla natura della luce

Il successo di questa teoria sembra vacillare nel momento in cui scopriamo l'**effetto fotoelettrico** e le sue conseguenze, che portano a concludere che la luce,

contrariamente a quanto ci è sembrato finora, sia costituita di corpuscoli e non abbia affatto un carattere ondulatorio. L'apparente paradosso (abbiamo esperimenti che dimostrano che la luce è un'onda ed esperimenti che dimostrano il contrario) si supera grazie alla **Meccanica Quantistica** che, insieme alla **teoria della relatività** rappresenta un pilastro fondamentale della fisica. Non solo per le innumerevoli (ancorché ignorate) applicazioni, ma anche perché con esse si chiarisce definitivamente il significato del **metodo sperimentale**: quel che si misura è, ed esiste, e non ha alcun senso chiedersi qual è il reale aspetto dell'Universo, se per questo s'intende qualcosa di non misurabile sperimentalmente. Quando misuriamo il tempo con un orologio a bordo di un satellite GPS, che marcia in modo diverso da uno a Terra, non ha senso chiedersi se sia l'orologio a marciare diversamente, il tempo a scorrere in maniera diversa o se semplicemente il tempo scorra nello stesso modo ovunque, ma noi non siamo in grado di *leggere* correttamente quell'orologio: se non esiste una maniera oggettiva di misurare il tempo indipendentemente dalla presenza di un orologio dobbiamo accettare il risultato sperimentale secondo il quale il tempo è relativo all'osservatore. È quanto si dice nell'introduzione a questo capitolo (e non è sempre stato così: per buona parte della storia umana, da **Parmenide** ad **Einstein**, ciò che è, è sempre stato indipendente da ciò che si può misurare).

È l'Universo a scegliere come deve funzionare e non saremo certo noi a imporre all'Universo la maniera in cui è tenuto a funzionare, solo perché non siamo capaci di comprenderne il funzionamento effettivo secondo le nostre convinzioni.

1.12 L'entropia e il secondo principio

Ma parlando di tempo è naturale chiedersi perché scorra sempre in avanti e mai all'indietro: cosa impedisce agli oggetti che osserviamo di comportarsi come se andassero all'indietro nel tempo? Perché il gas contenuto in una lattina di una bibita gassata

esce sempre con violenza quando la si apre e non accade mai che vi rientri o che dell'aria entri nella lattina? In fondo, secondo tutte le Leggi fisiche scoperte finora, questo dovrebbe essere possibile: non ci sono Leggi che lo vietano! La risposta a questa domanda verrà dal **secondo principio della termodinamica**, che ci porterà a considerare l'**entropia** come un'altra grandezza fisica di cui tenere conto.

1.13 Le particelle elementari

A questo punto sembra che il modello che ci siamo fatti della materia e delle forze cui è soggetta sia completo. Ma in realtà, dallo studio dettagliato delle forze di natura elettrica si capirà che il mondo dev'essere più complicato di così: le particelle che compongono la materia dell'Universo sono molte di più rispetto a quelle che possiamo ipotizzare studiando le interazioni che abbiamo studiato finora. **Scopriremo** che esistono intere famiglie di nuove particelle che interagiscono attraverso forze che si possono pensare, a loro volta, come l'effetto dello scambio di altre particelle.

Il ruolo della misura in Fisica

La **fisica** è una **scienza sperimentale**: questo significa che in fisica si possono solo trattare questioni che hanno carattere sperimentale, sulle quali cioè si possono fare **esperimenti**. Un esperimento non è, come pensano in molti, una prova, un tentativo. Un esperimento consiste nella **misura** di una qualche **grandezza fisica**.

Per poter definire un esperimento dunque è necessario definire bene cosa intendiamo per **grandezza fisica**. Potremmo dire che una grandezza fisica è una qualunque caratteristica di un qualche sistema che si può misurare. La **misurazione** di una grandezza fisica consiste nel processo che conduce alla **misura**, che ne è il risultato. Una grandezza fisica si definisce **operativamente** quando si descriva in modo non ambiguo la sequenza di azioni da compiere per giungere alla misura che non deve risultare soggettiva.

Osservate la Figura 2.1, nella quale si vede un'opera dell'artista **Marcel Duchamp**. Il fatto che quest'opera sia conservata in alcune repliche in diversi musei (una delle quali presso la **Galleria Nazionale d'Arte Moderna** a Roma, che vi consigliamo di visitare indipendentemente da quel che pensate dell'opera di Duchamp) indica che quella che si vede nella foto è considerata almeno da una parte delle persone un'opera d'arte. Di sicuro una parte del pubblico che visita la Galleria non è d'accordo con tale definizione. Lo stesso accade anche con opere o artisti meno controversi: le opere di **Picasso** sono per lo più giudicate come capolavori, ma ci sono persone che fanno fatica a vederci qualcosa d'interessante; anche **Modigliani** suscita sentimenti diversi. Lo stesso grado di arbitrarietà di giudizio spetta alla musica: **Wagner**, **Berio**, **Bartók** sono solo alcuni esempi di compositori la cui opera è giudicata



Figura 2.1 La fontana di **Marcel Duchamp**.

in maniera controversa. Non solo le opere d'arte si possono includere tra quelle il cui valore è soggettivo: il gusto delle pietanze lo è altrettanto (la **pajata**, tipico piatto romanesco, non è certamente tra quelli che i più considerano una prelibatezza, e lo stesso vale per caviale, ostriche, e quant'altro).

La bellezza di un'opera e persino il fatto che un oggetto si possa definire **opera d'arte** è in sostanza una questione sulla quale non ci può essere alcuna oggettività (anche se può esserci unanimità di

giudizio). Idem dicasi per la bontà del cibo o dell'intensità d'un sentimento come l'amore, la rabbia, l'invidia, etc.. Se ne conclude che queste qualità non possono essere oggetto di studio in fisica: semplicemente per la fisica **non esistono**. Naturalmente questo non vuol dire necessariamente che quelle qualità non esistano (sono certo di provare certi sentimenti o di apprezzare un buon piatto di bucatini all'amatriciana), ma che un fisico non può prendere in considerazione queste qualità come oggetto dei suoi studi. Per lo stesso motivo nessun fisico onesto potrebbe includere tra l'oggetto delle sue ricerche l'esistenza di Dio, la cui presenza non è rivelabile in maniera oggettiva e certa. Naturalmente questo non esclude che lo stesso fisico possa credere in Dio. Semplicemente non deve farsi guidare, nel suo lavoro, da convinzioni di carattere personale che non siano scientificamente dimostrabili.

In definitiva, in fisica, possiamo prendere in considerazione solo grandezze fisiche **definite operativamente**, il cui valore sia determinabile in maniera chiara e precisa attraverso una ben determinata sequenza di operazioni di misura. Quello che si fa in fisica è misurare grandezze fisiche e stabilire relazioni tra di esse sotto forma di equazioni che chiamiamo **leggi fisiche**. Una legge fisica non è in alcun modo simile a una legge ordinaria varata da un Parlamento: la prima rappresenta soltanto una relazione tra grandezze fisiche che è stata stabilita **sperimentalmente**, cioè attraverso la misurazione ripetuta e sistematica di almeno due grandezze fisiche; la seconda è il risultato di una decisione più o meno arbitraria. Quando si dice che **il moto dei pianeti obbedisce alle Leggi di Keplero** s'intende dire che l'osservazione del moto dei pianeti ha condotto i fisici (Keplero, nella fattispecie) a ritenere che certe caratteristiche misurabili di questo moto rivelino una qualche relazione tra loro, che si può esprimere attraverso un'equazione. I pianeti, infatti, se ne infischiano delle leggi che noi formuliamo e non si sognano neppure di *obbedire* a qualche tipo di regola imposta da noi umani! Sarebbe, in un certo senso, più corretto affermare che non sono i pianeti a obbedire alle Leggi di Keplero, ma che sono queste ultime a *obbedire* al moto dei pianeti.

2.1 Misure e teorie

Una volta eseguite le misure e determinate le relazioni esistenti tra le grandezze fisiche oggetto di studio, si possono formulare **teorie** a partire da queste. Una teoria è un insieme organico di equazioni che discendono matematicamente da una o più leggi fisiche determinate sperimentalmente, eventualmente attraverso l'assunzione (talvolta implicita) di ipotesi relative alla maniera di interpretare i risultati delle misure. Non esistono teorie fisiche che possano fare a meno di fatti sperimentali. In fisica, dunque, non esistono teorie **false**, come talvolta si legge su qualche libro o articolo, proprio perché le teorie non possono che essere il risultato dell'analisi di fatti sperimentali che non possono essere che **veri** (a meno che, evidentemente, non siano *inventati*, ma in questo caso si tratta di una frode). Perlomeno se all'aggettivo **falso** si dà il significato che comunemente spetta a questa parola.

Per chiarire il senso di quest'affermazione vale la pena fare un esempio, anche se non conoscete ancora le teorie di cui parliamo. Avrete però sicuramente sentito parlare della **teoria geocentrica** di **Tolomeo** secondo la quale la Terra si troverebbe al centro dell'Universo e il Sole e gli altri pianeti ruoterebbero attorno ad essa. La teoria fu superata da quella **eliocentrica** di **Copernico** che dimostrò come la teoria di Tolomeo fosse *falsa*, ponendo al centro dell'Universo il Sole. Quando una teoria fisica è dichiarata *falsa* s'intende dire in realtà che ne esiste un'altra che spiega gli stessi fenomeni in modo più semplice rispetto alla prima, dove più semplice significa con un minor numero di ipotesi non verificabili (o non verificate). Dal punto di vista di chi osserva il moto del Sole e degli altri corpi celesti stando sulla Terra, la teoria di Tolomeo è del tutto plausibile e quindi **vera** nel senso che, osservando il cielo, se ne ricava l'impressione che tutto ruoti attorno a noi. Se però si eseguono osservazioni più raffinate, si vede che certi pianeti sembrano muoversi sí su orbite circolari attorno alla Terra, ma in certi momenti sembrano tornare indietro rispetto alla posizione occupata nei giorni precedenti. Inizialmente quest'osservazione portò gli studiosi dell'epoca a formulare l'ipotesi

(aggiuntiva) che i pianeti si muovessero sempre su **orbite circolari** dette **epicicli**, il cui centro ruotava attorno alla Terra su un'orbita detta **deferente**. Quest'ipotesi permetteva di spiegare il cosiddetto **moto retrogrado** dei pianeti come osservato da Terra. La teoria originale di Tolomeo non poteva più essere *vera*, senza l'aggiunta di epicicli e deferenti, ma dev'essere considerata tale nel limite in cui le misure si eseguono con precisione modesta (fino a quando non ci si accorge, cioè, della presenza del moto retrogrado). Ben presto, anche la teoria di Tolomeo emendata attraverso l'inclusione di epicicli e deferenti, si dimostrò *falsa* nel senso che non riusciva più a spiegare il moto osservato dei corpi celesti che si discostava da quanto predetto da questa teoria. Ma fino a quando ci si limita a osservazioni (misure) grossolane la teoria è vera.

Copernico dimostrò che assumendo che il Sole si trovasse al centro dell'Universo, e non la Terra, si potevano spiegare le osservazioni senza ricorrere a ipotesi aggiuntive come quella dell'esistenza di epicicli e deferenti. Da questo punto di vista la teoria di Copernico rappresenta una semplificazione della teoria di Tolomeo.

Solo molto più tardi **Newton** dimostrò che i moti dei pianeti attorno al Sole si potevano spiegare ipotizzando che la forza che li muoveva era la stessa che provocava la caduta degli oggetti qui sulla Terra. Da questo punto di vista, la teoria di Newton, rappresenta un'ulteriore semplificazione rispetto a quella precedente perché diminuisce il numero di forze di cui bisogna ipotizzare l'esistenza per spiegare le osservazioni sperimentali. Questa teoria dimostrò presto di possedere quella che il filosofo **Imre Lakatos** chiama **euristica positiva**: la capacità cioè di spiegare fenomeni estranei a quelli considerati per giungere alla formulazione della teoria. La teoria di Newton, infatti, spiegava contemporaneamente il moto di tutti i corpi celesti e dei corpi soggetti alla forza di gravità sulla Terra, anche quando le misure si fecero più precise e si scoprirono alcune deviazioni rispetto alle previsioni della teoria. Per esempio, il moto della Terra attorno al Sole, predetto con la Teoria Newtoniana, non era esattamente quello osservato, ma includendo gli effetti dei corpi celesti più vicini

come la Luna, calcolabili grazie alla stessa teoria, si dimostrò che le misure si potevano spiegare perfettamente assumendo questa teoria come **vera**. Addirittura, quando si scoprì che il moto di Saturno non era esattamente identico a quello predetto usando la teoria, s'ipotizzò che esistesse almeno un altro corpo celeste che ne perturbava il moto sempre per mezzo della stessa forza: questo corpo celeste (poi battezzato con il nome di **Urano**) fu cercato e trovato esattamente dove doveva essere secondo la teoria di Newton.

Solo il moto di Mercurio rimase senza spiegazione fino all'inizio del secolo scorso: l'asse maggiore della sua orbita infatti, che ha la forma di un'ellisse, sembra ruotare attorno al Sole e questo fenomeno, noto col nome di **precessione del perielio** non si può spiegare con la teoria di Newton. Nei primi anni del secolo XX **Albert Einstein** formulò una teoria per spiegare un fatto sperimentale piuttosto curioso: la velocità della luce sembrava essere sempre la stessa, indipendentemente dallo stato di moto relativo rispetto a un osservatore. In altre parole, anche se si corre dietro a un raggio di luce alla sua stessa velocità, non lo si vede fermo, ma lo si vede sempre allontanarsi a 300 000 km/s. La formulazione di questa teoria portò a ritenere che la gravità fosse il risultato di una deformazione dello spazio la cui entità poteva essere predetta conoscendo semplicemente la massa degli oggetti. Quest'ipotesi consentiva di spiegare (in un modo che non è affatto facile da illustrare a questo livello) il moto di Mercurio! La teoria di Einstein, in sostanza, possiede anch'essa un'euristica positiva rispetto a quella di Newton e di conseguenza quest'ultima dev'essere considerata *falsa*. Questo tuttavia non implica che si debba abbandonare totalmente e in effetti la teoria di Newton si usa correntemente in astrofisica e in astronautica, semplicemente perché entro certi limiti funziona perfettamente (e non potrebbe essere altrimenti, derivando da fatti sperimentali).

L'aggettivo **falso**, in definitiva, in fisica non ha lo stesso significato che nel linguaggio comune.

2.2 Le misure di base

Il più semplice processo per la misurazione di una grandezza fisica consiste nel confronto della grandezza fisica da misurare con una ad essa omogenea (cioè della stessa natura). Una **lunghezza**, per esempio, è una grandezza fisica perché esiste una procedura molto semplice per misurarla. Scegliamo qualcosa che a nostro giudizio possieda questa caratteristica (un righello, una matita, un dito, ...) e lo assumiamo come **unità di misura**: un oggetto che possieda questa stessa caratteristica avrà una lunghezza pari al numero di volte (eventualmente anche non intero) che l'unità di misura entra nel campione da misurare. Naturalmente, data l'arbitrarietà con la quale si può scegliere l'unità di misura, lo stesso oggetto potrebbe assumere lunghezze diverse secondo l'unità adottata. Per questo, oltre al valore della misura, occorre riportare anche l'unità adoperata nel comunicare il risultato dell'operazione ad altri.

Anche la misura di **peso** si esegue per confronto con un'unità di misura¹.

Il **tempo** è la grandezza fisica che si misura con l'orologio. Non è necessario disporre di un vero e proprio orologio per misurarlo: prima della sua invenzione si poteva comunque definire una procedura per misurarlo basata sul confronto con la durata del giorno (che presenta alcuni problemi pratici che qui non discutiamo).

È chiaro che, per quanto arbitraria, la scelta dell'unità di misura dovrebbe seguire certi criteri che consentano di ottenere risultati stabili quando si misura la stessa grandezza: per esempio, le lunghezze si possono misurare in *dita*, scegliendo il diametro delle dita come unità di misura, ma se si chiede a qualcuno di versarci due dita di vino nel bicchiere, il risultato può essere molto diverso, secondo i casi.

¹I fisici usano spesso pignoleggiare sul significato della parola **peso**, che usano per indicare una forza, in relazione alla parola **massa** che usano per indicare quel che i comuni mortali chiamano peso. In questo frangente specifico la cosa non fa molta differenza perciò useremo la parola peso che comunque è proporzionale alla massa e in ogni caso si determina per confronto. Al momento comunque non abbiamo alcuna ragione per dover distinguere tra massa e peso.

Innanzitutto le dita non sono tutte uguali; anche se se ne sceglie uno, il diametro del dito scelto non è uniforme (e la sua sezione non ha nemmeno la forma di un cerchio per cui bisognerebbe specificare lungo quale asse lo prendiamo); anche in questo caso, il mio dito non ha lo stesso spessore di quello di un bambino o di quello di un altro qualunque essere umano! Non si tratta dunque di una scelta molto sensata. Detto questo, si tratta comunque di un problema pratico che si può risolvere in molti modi e che non presenta particolari problemi di natura fondamentale, perciò non vale la pena soffermarsi troppo su questo aspetto. La moderna definizione delle unità di misura impiegate in fisica (e non solo) è responsabilità del **Bureau International de Poids et Measures** (BIPM) che ha sede a Parigi, sulla cui pagina web si trovano le necessarie informazioni. In quello che si chiama il **Sistema Internazionale** o **SI** la lunghezza, la massa e il tempo sono grandezze fisiche definite come **fondamentali**, che si misurano cioè per confronto diretto. Le corrispondenti unità di misura sono il metro (m), il chilogrammo (kg) e il secondo (s, non sec come talvolta si legge).

Per ogni unità di misura si possono definire multipli e sottomultipli, che di solito sono espressi in base dieci. Del metro, ad esempio, si definiscono i sottomultipli **millimetro** (mm, corrispondente a 1/1000 di metro o 0.001 m), **centimetro** (cm, corrispondente a 1/100 di metro o 0.01 m) e **decimetro** (dm, 1/10 di metro o 0.1 m). Altri sottomultipli molto usati sono il **micrometro** o **micron**, pari a 10^{-6} m, che si indica con il simbolo μm e il **nanometro** (nm, corrispondente a 10^{-9} m), oltre all'Ångström, pari a 10^{-10} m, indicato con il simbolo Å. Per le lunghezze si adopera spesso il multiplo chiamato **chilometro** (km, pari a 1000 m). Per indicare il multiplo o il sottomultiplo si adopera un **prefisso**: uno o più caratteri che precedono il simbolo dell'unità di misura, che indicano la potenza di dieci per cui moltiplicare il numero indicato per esprimerlo nell'unità fondamentale.

I prefissi più usati sono riportati nella Tabella 2.1, dalla quale si evince che 1 cm corrisponde a 10^{-2} m, cioè a 1/100 di metro, e che si pronuncia *centi-metro*. Analogamente il milligrammo corrisponde a

prefisso	potenza	nome
Z	10^{21}	Zetta
E	10^{18}	Exa
P	10^{15}	Peta
T	10^{12}	Tera
G	10^9	Giga
M	10^6	Mega
k	10^3	kilo
h	10^2	etto
da	10	deca
d	10^{-1}	deci
c	10^{-2}	centi
m	10^{-3}	milli
μ	10^{-6}	micro
n	10^{-9}	nano
p	10^{-12}	pico
f	10^{-15}	femto
a	10^{-18}	atto
z	10^{-21}	zepto

Tavola 2.1 I prefissi usati nei multipli e sottomultipli delle unità di misura.

10^{-3} g, e il TW (Terawatt) corrisponde a 10^{12} W.

Dal momento che la scelta delle unità di misura da impiegare durante una misura è arbitraria è necessario saper esprimere una grandezza fisica nei diversi sistemi. La lunghezza di una matita, evidentemente, resta la stessa sia che la misuriamo in metri, che nel caso in cui la misuriamo in cm. Evidentemente la stessa matita deve avere la stessa lunghezza anche se misurata in **pollici**, un'unità che si usa nei Paesi anglosassoni. Per passare da un sistema di unità di misura A all'altro B è sufficiente conoscere a quante unità del sistema A corrisponde l'unità di misura nel sistema B . Si tratta quindi il simbolo che indica l'unità di misura come una quantità algebrica. Un esempio vale più di mille parole, in questo caso: supponiamo di possedere un televisore il cui schermo ha una diagonale che misura 56 cm. La lunghezza della diagonale dello schermo è quella che caratterizza le dimensioni di uno schermo televisivo (conoscendone la diagonale e sapendo che il

Multipli e prefissi in informatica

In informatica è più pratico usare le potenze di due invece che quelle di dieci. La memoria di un computer o lo spazio occupato da un *file* si misurano in byte (B). Perciò 1 kB corrisponde a 1024 B e non a 1000 B. 1024 infatti è una potenza di 2 ($1024 = 2^{10}$). Corrispondentemente cambiano, anche se di poco, i fattori indicati dai prefissi.

Il Megabyte (1 MB) corrisponde a circa un milione di byte, cioè a circa mille kB. Per la precisione si tratta di 1024 volte un kB, cioè 1048 576 B. Analogamente 1 GB = 1024 MB = 1073 741 824 B e 1 TB = 1024 GB = 1099 511 627 776 B.

Come si vede la differenza con i prefissi ordinari non è molto alta e dunque si usano spesso questi ultimi per stimare rozzamente le dimensioni effettive di una memoria.

rapporto tra la base e l'altezza è di 16 : 9, possiamo conoscerle tutte). Per comprarne uno uguale dobbiamo sapere quanto è lunga la diagonale in pollici. Sappiamo che la lunghezza di un centimetro corrisponde a circa 0.39 pollici (il cui simbolo è un doppio apice: "), cioè che

$$1 \text{ cm} = 0.39'' . \quad (2.1)$$

Trattando le unità di misura come fossero quantità algebriche possiamo trattare questa equivalenza come un'equazione e scrivere che

$$1 = \frac{0.39''}{\text{cm}} \quad (2.2)$$

così che la misura di 56 cm si potrebbe esprimere come $56 \times 1 \text{ cm}$ e, sostituendo a 1 il *valore* trovato sopra, ottenere

$$56 \text{ cm} = 56 \times 1 \text{ cm} = 56 \times \frac{0.39''}{\text{cm}} \text{ cm} = 21.84'' \quad (2.3)$$

Il televisore è da 22 pollici. Se disponete di un collegamento a Internet, è facile operare trasformazioni di unità di misura ricorrendo al motore di ricerca

Google: nel campo destinato alle parole da cercare digitate la misura completa del simbolo dell'unità, seguita dalla parola **to** seguita, a sua volta, dall'unità nella quale si vuole convertire (tutto in inglese: nel nostro caso avremmo dovuto scrivere **56 cm to inches**).

In certi casi la misura di una grandezza fisica si ottiene combinando più misure di altre grandezze fisiche. In questo caso la grandezza fisica in questione si dice **derivata**. Quando si combinano più misure per costruire una grandezza fisica derivata, l'unità di misura della grandezza fisica in esame si ottiene moltiplicando o dividendo tra loro le unità di misura delle grandezze fisiche che sono state usate per determinarla. Per conoscere l'area di un rettangolo si devono misurare le lunghezze dei suoi lati: se il lato del rettangolo si misura in metri (m), l'area si misura in metri quadri (m²) e si dice che ha le **dimensioni fisiche** di una lunghezza al quadrato. La natura della grandezza fisica si esprime scrivendo il simbolo corrispondente (L per le lunghezze) tra parentesi quadre, come in

$$[\text{lato del rettangolo}] = [L] \quad (2.4)$$

e in

$$[\text{area del rettangolo}] = [L^2] . \quad (2.5)$$

La velocità di un corpo che si muove rappresenta la distanza percorsa per unità di tempo e quindi si misura determinando quanto spazio è stato percorso in un'unità di tempo e si esprime come il rapporto tra una distanza e un tempo. Di conseguenza ha le dimensioni di una lunghezza divisa per un tempo, cioè

$$[\text{velocità}] = [LT^{-1}] . \quad (2.6)$$

e si misura in m/s o in multipli di queste unità (per esempio in km/h).

È utile osservare che la natura fondamentale di una grandezza fisica è anch'essa il risultato di una scelta arbitraria. Ad esempio, un'altra possibile scelta consiste nel definire le velocità come grandezze fisiche fondamentali che si misurano confrontandole

con quella della luce. Un corpo che nel SI si muove alla velocità di 200 000 km/s, in questo sistema si muove alla velocità di $2/3$. La velocità quindi si può considerare **adimensionale**, cioè priva di **dimensioni fisiche**, nel senso che la sua misura è sempre esprimibile come un rapporto con una grandezza fisica omogenea il cui valore è fissato dalla Natura. Anche nel SI, la lunghezza si esprime come un rapporto tra quella dell'oggetto da misurare e quella del metro campione, ma la lunghezza di quest'ultimo è il risultato di una nostra scelta e perciò si dice che la lunghezza **ha** dimensioni fisiche (quelle di una lunghezza, appunto). Se le velocità si misurano in frazioni di velocità della luce e i tempi si misurano in secondi, le lunghezze non hanno più un carattere fondamentale, ma sono grandezze fisiche **derivate**: si ricavano cioè dalla combinazione di misure di natura diversa. Per sapere quanto è lungo un oggetto si deve misurare il tempo che impiega la luce per giungere da un estremo all'altro: il valore della misura dunque si esprime in unità di velocità della luce (che è adimensionale), moltiplicate per le unità di misura del tempo (s). In questo sistema dunque le lunghezze hanno le stesse dimensioni fisiche dei tempi e si misurano in secondi².

Dal momento che una legge fisica è una relazione tra grandezze fisiche che si esprime come un'equazione o una disequazione, è evidente che una grandezza fisica può solo essere uguale (o comparata) a un'altra grandezza fisica ad essa omogenea, perciò è impossibile che in una legge fisica compaiano a primo membro grandezze fisiche con dimensioni diverse da quelle che compaiono a secondo membro. I due membri della relazione devono necessariamente avere le stesse dimensioni fisiche realizzando quella che si chiama un'**equazione dimensionale**. La legge fisica che esprime lo spazio percorso in un tempo t da un corpo che si muove a velocità costante pari a v è

$$x = x_0 + vt \quad (2.7)$$

²L'anno luce, in effetti, è una misura di lunghezza corrispondente alla distanza che la luce percorre in un anno. In questo caso un'unità di misura di tempo (l'anno) s'impiega per indicare una distanza.

dove x_0 rappresenta la posizione iniziale del corpo, espressa come la distanza da un punto scelto come riferimento. A primo membro le dimensioni di x sono quelle di una lunghezza $[x] = [L]$. A secondo membro quindi dev'esserci una lunghezza perché una lunghezza può solo essere uguale a una lunghezza! Il secondo membro è formato dalla somma di due termini, ciascuno dei quali dev'essere una lunghezza e per lo stesso motivo le somme o le sottrazioni tra grandezze fisiche si possono eseguire solo tra grandezze tra loro omogenee. x_0 è una lunghezza $[x_0] = [L]$, quindi resta da controllare che lo sia il prodotto vt . v è una velocità che ha le dimensioni di una lunghezza diviso un tempo: $[v] = [LT^{-1}]$. Quando la si moltiplica per un tempo le dimensioni fisiche del prodotto sono

$$[LT^{-1}][T] = [LT^{-1}T] = [L], \quad (2.8)$$

come ci si aspetta se l'equazione è corretta. Le equazioni dimensionali sono un formidabile strumento di controllo della correttezza di un risultato e bisogna imparare a servirsene. Sebbene nessuno possa garantire che un'equazione dimensionalmente corretta sia giusta, è invece certo che un'equazione dimensionalmente incoerente è senza alcun dubbio sbagliata! Se scrivessimo, per errore, che

$$x = x_0 + vt^2 \quad (2.9)$$

vedremmo subito che il secondo addendo a secondo membro ha le dimensioni di una velocità per un tempo al quadrato, cioè di $[LT^{-1}T^2] = [LT]$, che non sono quelle giuste, perché a primo membro c'è una lunghezza!

2.3 Gli strumenti

Il modo più semplice di eseguire una misura consiste nel confrontare direttamente l'oggetto da misurare con uno **strumento**: una copia del campione dell'unità di misura che dunque possiede le stesse caratteristiche dell'oggetto da misurare.

Per esempio, per misurare un volume, potete usare un bicchiere da birra da una **pinta**, che corrisponde a quasi mezzo litro. Con questo strumento potete



Figura 2.2 Il boccale di birra da una pinta è uno strumento di misura diretta di volumi, che contiene 0.473 l di liquido. La caraffa invece è uno strumento graduato, perché riporta l'indicazione di frazioni dell'intero.

misurare volumi pari a un numero intero di pinte. Non potete misurare frazioni di questo volume, perché la forma del bicchiere non è regolare e l'assenza di tacche che riportino i valori intermedi impedisce di fare una misura (si può fare una stima, non una misura).

Se invece di un boccale da birra si usa una caraffa graduata, si possono misurare anche volumi diversi da quello della caraffa intera. Lo strumento si dice **graduato**.

Se dovete preparare le tagliatelle fatte in casa dovete usare una certa quantità di farina, tipicamente espressa in grammi. Se non avete una bilancia potete misurare il volume della farina e, conoscendone la densità, ricavare il valore che vi serve attraverso un'operazione di **misura indiretta**: in pratica non si misura direttamente la grandezza fisica cui si è interessati, ma una o più grandezze che si sa possedere una relazione con quella principale. Sapendo, ad esempio, che un volume di farina pari a un litro pesa circa 800 g, avendo bisogno di 500 g di farina servono

$$V = 500 \text{ g} \times \frac{1 \text{ l}}{800 \text{ g}} = 0.625 \text{ l} \quad (2.10)$$

di farina. Invece di dover fare ogni volta questi conti, potete comprare in un negozio di casalinghi una ca-



Figura 2.3 Questa caraffa è uno strumento tarato: di fatto esegue una misura di volume, che è convertita in una misura di massa grazie alla conoscenza della densità della farina.

raffa graduata sulla quale sono riportate, oltre alle indicazioni relative al volume, anche quelle relative alla massa di farina e zucchero. Questo strumento si dice **tarato** perché riporta sulla propria scala i valori di una grandezza fisica (la massa) non **omogenea** a quella che effettivamente misura (il volume). La scala si ottiene attraverso un'operazione di **taratura** o **calibrazione**, consistente nello stabilire la relazione esistente tra la grandezza fisica di cui fornire la misura e quella effettivamente misurata.

Gli strumenti non sono tutti uguali. Anche per la misura della stessa grandezza fisica, sono necessari strumenti con proprietà diverse, secondo le necessità. La misura della larghezza di un marciapiede richiede strumenti diversi da quelli necessari per la misura del diametro di un capello.

Una delle principali caratteristiche di uno strumento è la sua **sensibilità**, che rappresenta la più piccola variazione di una grandezza fisica che lo strumento è in grado di apprezzare o rivelare.

La sensibilità di uno strumento è spesso confusa con la sua **precisione**, che è una caratteristica

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 2.4 La sensibilità limitata di questa bilancia non permette di misurare la variazione di peso conseguente all'aggiunta di una banconota agli oggetti sul piatto [https://www.youtube.com/watch?v=XKbn0z71aUQ]. Per questo strumento il peso della banconota è di fatto nullo.

che indica quanto due misure della stessa grandezza, eseguite con lo stesso strumento, si discostano. Minore è lo scostamento, maggiore è la precisione dello strumento. La precisione può dipendere dalle condizioni in cui viene eseguita la misura. Paradossalmente, uno strumento molto preciso, in generale è anche poco sensibile e viceversa. Infatti alta precisione significa risultati sempre uguali, il che è possibile per strumenti poco sensibili. L'**accuratezza** di uno strumento, invece, ne rappresenta la fedeltà rispetto al campione. Un metro è tanto più accurato quanto più la sua lunghezza coincide con quella del campione ufficiale. Noi stessi potremmo realizzare un metro ritagliando una striscia di carta di lunghezza opportuna. Quest'ultimo strumento potrebbe facilmente essere poco accurato perché risulterebbe probabilmente un po' più grande o un po' più piccolo di un vero metro. La sua sensibilità e la sua precisione sono indipendenti dall'accuratezza.

Un'ultima importante caratteristica di uno strumento è la sua **portata** o **fondo scala** che indica il massimo valore di una grandezza fisica che lo strumento è in grado di misurare. Per un metro da muratore la portata è tipicamente di due metri. La portata di un righello è di qualche decina di centimetri, mentre il metro a nastro usato dai geometri può avere una portata di venti metri (doppio decametro).

Esercizio 2.1 *Caratteristiche degli strumenti*

Cerca almeno cinque strumenti di misura in casa (righello, bilancia, orologio, termometro, goniometro, cilindro graduato, etc.). Identifica il tipo di strumento e stima, se possibile, le caratteristiche di portata, precisione, accuratezza e sensibilità. Esegui alcune misure con ciascuno di essi. Se possiedi uno strumento con una doppia scala (per esempio un righello con indicate le lunghezze in cm e in pollici, una bilancia con indicazione in kg e libbre) esegui alcune misure e fai un grafico del valore in un'unità in funzione del valore nell'altra unità. Prova a ricavare, da questo grafico, la relazione che esiste tra le diverse unità di misura.

2.4 Notazione scientifica

Risulta molto comodo esprimere i multipli di dieci come potenze nella così detta **notazione scientifica** che consiste nell'esprimere il valore di una misura come un numero compreso tra 1 e 10, moltiplicato per un'opportuna potenza di dieci. La distanza tra il Sole e la Terra L_{ST} , ad esempio, è di 149 milioni di km, circa. Se dovessimo esprimere questo numero in unità del SI, dovremmo scrivere

$$L_{ST} = 149\,000\,000\,000\,000\text{ m}, \quad (2.11)$$

ricordando che 1 km=1000 m. La stessa distanza si può esprimere come

$$L_{ST} = 149 \times 10^{12} \text{ m} = 1.49 \times 10^{14} \text{ m}. \quad (2.12)$$

Per passare alla notazione scientifica abbiamo prima contato il numero di zeri (12) scrivendo L_{ST} come un numero, detto **mantissa** moltiplicato per dieci elevato al numero di zeri trovati. Quindi abbiamo diviso la mantissa per 100, in modo da scriverla come un numero più piccolo di dieci. Il numero di cifre dopo la virgola (due in questo caso) ci dice la potenza di dieci per cui occorre moltiplicare il numero trovato per riottenere quello originale. Sfruttando

quindi le proprietà delle potenze abbiamo scritto $10^2 \times 10^{12} = 10^{14}$.

La notazione non è solo più compatta. Permette anche di eseguire i calcoli rapidamente in maniera semplificata o approssimata. Se, ad esempio, volessimo calcolare la distanza tra il Sole e la Luna sapendo che quest'ultima dista dalla Terra $L_{TL} = 300\,000 \text{ km}$, dovremmo sommare tra loro due numeri molto grandi e poco maneggevoli. Se però esprimiamo L_{TL} in notazione scientifica come $L_{TL} = 3 \times 10^8 \text{ m}$ vediamo subito che l'esponente di 10 nei due casi è molto diverso. Quello di L_{TL} , 8, è parecchio più piccolo di 14, che è quello di L_{ST} (ricordate che ogni unità corrisponde a una differenza di un fattore 10). L_{TL} è del tutto trascurabile rispetto a L_{ST} e possiamo dire che $L_{ST} + L_{TL} \simeq L_{ST}$. Se invece dobbiamo sommare a L_{ST} una distanza pari a $d = 324.26$ milioni di km, possiamo scrivere d in notazione scientifica come $d = 3.2426 \times 10^{14} \text{ m}$ ed eseguire la somma usando la proprietà associativa:

$$\begin{aligned} L_{ST} + d &= (1.49 + 3.2426) \times 10^{14} \text{ m} \\ &= 4.7326 \times 10^{14} \text{ m}. \end{aligned} \quad (2.13)$$

L'uso della notazione scientifica rende anche evidente l'**ordine di grandezza** di una misura, che ne definisce la **scala** alla quale avviene un determinato fenomeno. L'ordine di grandezza è rappresentato dall'esponente del valore della misura espressa in notazione scientifica. Per esempio, l'ordine di grandezza della distanza Sole-Terra è 10^{14} m , mentre quello della distanza Terra-Luna è di 10^5 km (notate che abbiamo espresso gli ordini di grandezza usando due diverse unità di misura). Come abbiamo visto, conoscendo l'ordine di grandezza, si può immediatamente stabilire se due grandezze tra loro omogenee sono comparabili oppure no (purché siano espresse nelle stesse unità): la distanza tra il Sole e la Terra è enormemente maggiore di quella tra Terra e Luna e la seconda è trascurabile rispetto alla prima.

Avrete notato che, prima di eseguire la somma, abbiamo fatto in modo di esprimere tutte le quantità secondo grandezze tra loro identiche. Non possiamo sommare una lunghezza espressa in metri

con una espressa in km. Prima dobbiamo avere le due grandezze espresse nella stessa unità di misura. Il risultato sarà la grandezza espressa nelle unità impiegate.

2.5 Un esperimento istruttivo

Fate il seguente esperimento: appendete un peso a un filo e fatelo oscillare, misurando con un cronometro il periodo di oscillazione (cioè il tempo che intercorre tra quando il peso si trova in una certa condizione, per esempio quella di partenza, e l'istante in cui torna a occupare la stessa condizione). Non serve disporre di strumentazione molto sofisticata: basta prendere un mazzo di chiavi appeso a un laccetto e usare uno smartphone come cronometro (ci sono centinaia di *app* gratuite che consentono di misurare tempi con precisioni del centesimo di secondo).

Eseguite la misura insieme a un compagno o una compagna (se le prime volte ci sono problemi nello stabilire l'istante di partenza o nel manovrare i cronometri ripetete la misura fino a quando non vi sentite sicuri). Molto probabilmente troverete valori diversi: poniamo che siano 1.13 s e 1.3 s. Che succede? Il tempo di oscillazione delle chiavi dev'essere lo stesso, indipendentemente da chi lo misura. Uno di due deve aver sbagliato. Ripetiamo l'operazione. Al secondo tentativo nessuna delle due nuove misure è uguale a quella di prima (1.1 s e 1.25 s nell'esercizio che hanno svolto i nostri studenti). Vuol dire che si sono sbagliati tutti? Come si fa a decidere qual è la misura giusta e qual è quella sbagliata? Continuiamo a misurare, ottenendo valori sempre diversi, solo talvolta uguali a quelli già ottenuti. Nella Tavola 2.2 sono riportate le misure eseguite da due studenti (Anna e Bruno).

Prima di analizzare il contenuto della tabella osserviamo il modo in cui è stata redatta: avrete notato che il primo tempo misurato da Bruno di 1.3 s in tabella è stato rappresentato come 1.30 s, aggiungendo uno zero non significativo. Quello zero, in effetti, non è **matematicamente** significativo, ma lo è invece per la fisica. Indica che lo strumento ado-

n	$T_A(s)$	$T_B(s)$	n	$T_A(s)$	$T_B(s)$
1	1.13	1.30	14	1.06	1.35
2	1.10	1.25	15	1.05	1.35
3	1.02	1.49	16	1.02	1.34
4	1.09	1.40	17	1.14	1.45
5	1.10	1.43	18	1.14	1.31
6	1.10	1.36	19	1.14	1.27
7	1.09	1.39	20	1.14	1.35
8	1.06	1.40	21	1.13	1.35
9	1.06	1.34	22	1.21	1.28
10	1.10	1.25	23	1.17	1.28
11	1.14	1.04	24	1.25	1.41
12	1.14	1.41	25	1.14	1.35
13	1.18	1.32	26	1.13	1.40

Tavola 2.2 Misure eseguite da due studenti del periodo di oscillazione di un pendolo.

perato da Bruno ha una sensibilità dell'ordine del centesimo di secondo. La misura eseguita da Bruno quindi è 1.30 s perché Bruno ha potuto leggere sul display del suo smartphone il numero 1.30 e non 1.3.

Anna e Bruno hanno eseguito 26 misurazioni e il fatto che le misure non siano sempre uguali l'una all'altra indica che nel processo di misura sono presenti elementi che fanno **fluttuare** i valori in un modo che sembra casuale. In effetti è facile interpretare il risultato ottenuto assumendo che effettivamente il mazzo di chiavi oscilli sempre con lo stesso periodo, ma che ogni volta che si esegue la misura i riflessi degli studenti non sono costanti e talvolta gli studenti fanno partire il cronometro un po' in anticipo, talvolta un po' in ritardo (e lo stesso accade quando lo fermano). Se sostituissimo gli studenti con fotocellule o con altri sistemi (in questi casi è utile un'*app* chiamata **Physics Gizmo**, disponibile per smartphone Android, che permette di usare il sensore di prossimità dei telefoni per far partire e fermare un timer) otterremmo probabilmente una minore variabilità, che tuttavia non sarebbe annullata completamente. Il fatto è che a ogni tentativo anche le condizioni del mazzo di chiavi non sono le stesse: non è detto che parta sempre rigorosamente

da fermo: talvolta lo sperimentatore applica involontariamente una piccola spinta. Anche la resistenza opposta dall'aria al moto delle chiavi influisce su questo tempo. Questa resistenza a sua volta dipende dalle condizioni locali dell'aria come velocità, umidità, temperatura, pressione, etc.. Insomma ci sono decine di effetti che possono contribuire ad aggiungere o a sottrarre al tempo *vero* T una quantità di tempo variabile di volta in volta. Insomma, il risultato di una misura non è un valore certo, unico, perfettamente determinato: è sempre il risultato del sommarsi di tanti effetti casuali e pertanto è esso stesso un **numero casuale**. Questo non vuol dire affatto che è del tutto privo di significato!

2.6 Proprietà statistiche delle variabili casuali

Per capire come gli effetti casuali influenzino le misure dobbiamo studiare il comportamento delle variabili casuali. Dal momento che questo è un manuale di fisica sperimentale e non di matematica, per capire le proprietà delle variabili casuali, possiamo fare qualche esperimento con qualcosa che abbia natura **aleatoria**, come il lancio di un dado, lasciando una trattazione formale e rigorosa dell'argomento all'insegnante di matematica.

Lanciando N volte un dado si possono ottenere tutti i possibili punteggi da 1 a 6. Se si conta il numero $n(x)$ di volte che esce il punteggio x e si riporta su un grafico $n(x)$ in funzione del punteggio stesso si costruisce quel che si chiama un **istogramma**. Indicando con $\{k\}$ l'insieme di tutti i possibili valori della variabile casuale x e con k uno dei suoi elementi, costruendo un istogramma, per ciascuno dei possibili valori della variabile casuale $x \in \{k\}$, si riporta il numero di volte $n(x)$ che $x = k$ per ciascuno dei possibili valori che k può assumere. Si può anche fare un istogramma riportando in funzione dei valori possibili della variabile casuale, questo stesso numero **normalizzato**, cioè diviso per il numero totale di estrazioni N , $f(x) = n(x)/N$ che prende il nome di **frequenza**.

La somma di tutti gli $n(k)$, per ognuno dei valori possibili k , dev'essere evidentemente uguale a N :

$$\sum_{\{k\}} n(k) = N \quad (2.14)$$

e quella delle $f(x)$ dev'essere naturalmente uguale a 1:

$$\sum_{\{k\}} f(k) = \sum_{\{k\}} \frac{n(k)}{N} = \frac{1}{N} \sum_{\{k\}} n(k) = \frac{N}{N} = 1. \quad (2.15)$$

Di conseguenza, tenendo conto del fatto che $n(k) \geq 0$, dev'essere sempre $0 \leq f(k) \leq 1$. Le condizioni (2.14) e (2.15) si chiamano **condizioni di normalizzazione**.

Fate l'esercizio di lanciare 100 volte un dado e scrivete su un foglio elettronico i punteggi ottenuti³. Nell'esercizio che abbiamo fatto noi i primi dieci numeri estratti sono

$$\{x\} = 5, 1, 5, 5, 1, 6, 2, 2, 5, 2. \quad (2.16)$$

Costruiamo l'istogramma contando il numero di volte in cui compare ciascuno dei possibili valori $k = 1, \dots, 6$: $n(1) = 2$, $n(2) = 3$, $n(3) = n(4) = 0$, $n(5) = 4$ e $n(6) = 1$. Corrispondentemente le frequenze sono $f(1) = 0.2$, $f(2) = 0.3$, $f(3) = f(4) = 0$, $f(5) = 0.4$ e $f(6) = 0.1$. La variabilità è piuttosto ampia e non sembrano esserci regole precise seguite da questi numeri. Provando con 30 lanci i valori che abbiamo trovato sono riportati in Tabella 2.3 (non vi fidate e fate da soli l'esperimento).

Si può notare come, all'aumentare di N , le $n(x)$ assumono valori sempre più simili tra loro e, di con-

³Invece di lanciare un vero dado, per far prima potete contare sulla capacità dei computer di estrarre numeri casuali con le stesse proprietà statistiche dei punteggi di un dado. Basta scrivere l'opportuna funzione in ciascuna cella del foglio per ottenere un numero casuale. Secondo i casi (Excel, Numbers, OpenOffice, Google Spreadsheet, etc.) le funzioni hanno nomi diversi, ma il loro nome comincia quasi sempre pre RAND. Per esempio, usando i fogli elettronici di Google potete usare la funzione `RANDBETWEEN(1,6)` per estrarre numeri casuali compresi tra 1 e 6.

	$N = 10$		$N = 30$		$N = 100$		$N = 1000$	
k	$n(k)$	$f(k)$	$n(k)$	$f(k)$	$n(k)$	$f(k)$	$n(k)$	$f(k)$
1	2	0.2	3	0.10	15	0.15	162	0.162
2	3	0.3	8	0.27	17	0.17	178	0.178
3	0	0.0	2	0.07	16	0.16	165	0.165
4	0	0.0	3	0.10	17	0.17	173	0.173
5	4	0.4	8	0.27	20	0.20	170	0.170
6	1	0.1	6	0.20	15	0.15	152	0.152

Tavola 2.3 Istogramma dei valori e delle frequenze per i lanci di un dado eseguito $N = 10$, $N = 30$ e $N = 100$ volte.

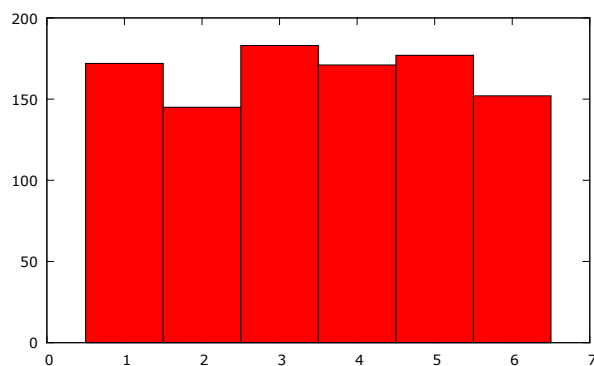


Figura 2.5 Distribuzione dei punteggi relativi a 1000 lanci di un dado.

tutto lo spazio dei possibili valori di k . Quello che si vede sperimentalmente è che i valori della frequenza $f(k)$, all'aumentare di N diventano via via più uniformi e sempre più simili proprio a $1/6 \simeq 0.167$. Possiamo quindi dire, solo sulla base dell'esperienza, che

$$f(k) \simeq p(k) \quad (2.17)$$

e che quest'affermazione è tanto più vera quanto maggiore è il numero N di prove effettuate. Ci si può quindi attendere che, se fosse possibile fare infinite prove si dovrebbe avere $f(k) = p(k)$. Questo si scrive come

$$\lim_{N \rightarrow \infty} f(k) = p(k). \quad (2.18)$$

sequenza, le $f(x)$ tendono via via ad assumere valori sempre più vicini a 0.17. In particolare, se per $n = 10$ la differenza tra il valore massimo e minimo di $f(x)$ è 0.4, per $N = 30$ diventa 0.20 fino ad arrivare, per $N = 100$ a 0.05. Addirittura per $N = 1000$ (la cui distribuzione di può vedere in Fig. 2.5) questa differenza si riduce a 0.026.

Se definiamo la **probabilità** di ottenere il punteggio k , $p(k)$, come il rapporto tra il numero dei casi in cui si può ottenere k (1) e quello di tutti i casi possibili (6), abbiamo che $p(k) = 1/6 \simeq 0.167$ per ogni valore di k . Diciamo quindi che i risultati del lancio di un dado sono distribuiti **uniformemente** nel senso che la probabilità di ottenere uno qualunque dei punteggi è uniforme, cioè la stessa, in

Questo è uno dei modi in cui si esprime quella che si chiama **Legge dei grandi numeri**.

Se invece di lanciare un solo dado ne lanciamo due, i possibili punteggi che possiamo ottenere vanno da 2 a 12, ma, se ci pensate un attimo, la distribuzione di questi punteggi non è più uniforme (potete vederla nella Fig. 2.6). Il valore 2 si può ottenere, così come il valore 12, solo se tutti e due i dadi presentano lo stesso punteggio: 1 + 1 oppure 6 + 6. Il punteggio 7, invece, si può ottenere in una moltitudine di modi: 1 + 6, 2 + 5, 3 + 4. Questo implica che il punteggio 7 ha una probabilità che è un fattore 3 più alta rispetto a quella di ottenere i

La statistica di Bayes

La statistica moderna considera inadeguata la definizione di probabilità data nel testo. In effetti, a essere rigorosi, per poter dire che la probabilità di uscita di uno dei punteggi del dado è di $1/6$, è necessario ammettere che tutti i punteggi hanno uguale probabilità: il problema quindi è che per definire una probabilità occorre usare il concetto stesso di probabilità. A metà del 1700 **Thomas Bayes** rese esplicito il problema sostenendo che fosse impossibile (ma in fondo inutile) definire una **probabilità oggettiva** e propose di adottare una statistica nella quale la probabilità avesse un carattere **soggettivo**, dipendente cioè da quanto ciascuno confidasse nel verificarsi di un evento.

Per quanto possa sembrare strano, l'adozione di una probabilità soggettiva non impedisce di fare previsioni anche piuttosto precise e la teoria di Bayes conduce a risultati verificabili. La moderna statistica si basa dunque sui risultati della teoria Bayesiana e non sull'approccio **frequentista** adottato nel testo. Va detto, tuttavia, che per la maggior parte dei casi le due teorie, quella di Bayes e quella frequentista, conducono a risultati coincidenti. Lasciamo perciò la trattazione formalmente corretta all'insegnante di matematica, adottando l'approccio classico perché dal punto di vista formale è più semplice.

valori 2 o 12. La distribuzione non è più uniforme, ma presenta un picco attorno al valore 7.

Se si sommano i punteggi di tre dadi, i punteggi totali risultano piccati attorno al valore 11, mentre se se ne sommano quattro il risultato è una distribuzione ancora più piccata attorno al valore 14 (fate sempre qualche prova).

In generale, sommando i punteggi di m dadi si ottengono distribuzioni che sono sempre più piccate su un valore che si può prevedere essere uguale a $m \times 3.5$, dove 3.5 è il punto centrale della distribuzione (uniforme) dei punteggi dei singoli dadi. Al crescere del numero di dadi lanciati contemporanea-

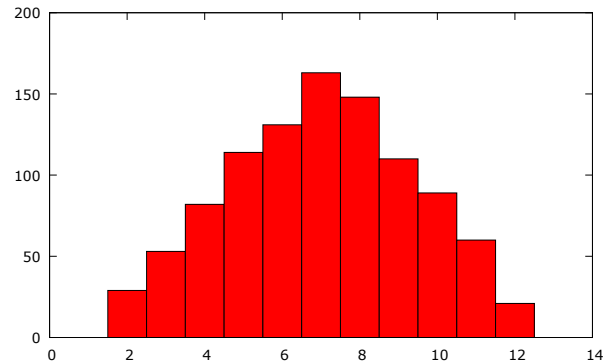


Figura 2.6 Distribuzione dei punteggi relativi a 1000 lanci di due dadi.

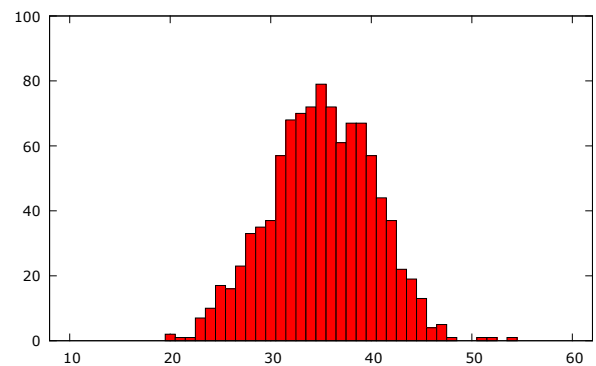


Figura 2.7 Distribuzione dei punteggi relativi a 1000 lanci di dieci dadi.

mente, la distribuzione assume sempre di più una forma *a campana* simmetrica. Nella Fig. 2.7 si vede la distribuzione del punteggio ottenuto lanciando dieci dadi.

Possiamo aspettarci qualcosa di simile anche per altri fenomeni a carattere aleatorio. Se per esempio misurate l'altezza o il peso di ciascuno studente in una classe, troverete una distribuzione più o meno simmetrica e più o meno stretta. Se sommate tutte le altezze o i pesi degli studenti di una classe e ripetete l'esperimento sugli studenti di altre classi, troverete che la somma delle altezze e la somma dei pesi sono distribuite anch'esse in modo da realizzare una distribuzione a campana, del tutto simile a quella che si ottiene con i dadi.

In generale, quello che si scopre facendo molti esperimenti del genere è che, qualunque sia la distribuzione delle variabili casuali iniziali, la loro somma si distribuisce sempre come una campana, come del resto si distribuiscono i dati relativi ai tempi di oscillazione del mazzo di chiavi raccolti da Anna e Bruno di cui si parla al paragrafo precedente.

Questa è una proprietà generale delle variabili casuali, la cui somma, indipendentemente dalla forma della distribuzione della singola variabile, tende a distribuirsi, al crescere del numero di addendi, in modo tale da assumere la forma di quella che si chiama una **gaussiana** o **curva di Gauss** che si rappresenta matematicamente come una funzione

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right). \quad (2.19)$$

La massima frequenza si ha per $x = \mu$ (in questo caso l'argomento dell'esponenziale vale 0 e l'esponenziale quindi vale 1, per cui $f(x) = 1/\sqrt{2\pi}\sigma$). La massima frequenza dipende da σ che in un certo senso ci dice quanto è larga la curva: più è grande σ più la curva assume una forma allargata, mentre per σ piccoli la curva è stretta. Quando $x = \mu \pm \sigma$ la curva assume lo stesso valore che è pari a

$$f(x) = \frac{1}{\sqrt{2e\pi}\sigma}, \quad (2.20)$$

perciò, rispetto al suo valore massimo, si riduce a un fattore $1/\sqrt{e} \simeq 0.6$. In Fig. 2.8 si vede una gaussiana i cui parametri sono $\mu = 35$ e $\sigma = 5.48$ (i valori attesi per il lancio di dieci dadi).

Eseguendo numerose misure della stessa grandezza fisica quello che si trova è che i valori si distribuiscono praticamente sempre come una gaussiana (a meno che non siano tutti uguali). Dal momento che la curva di Gauss è quella che descrive la distribuzione della somma di variabili casuali qualunque, possiamo pensare che il risultato di una misura sia la conseguenza del sommarsi di numerosi effetti che influenzano il risultato della misura. Qualunque sia la distribuzione di probabilità che si verifichi un determinato effetto, il risultato sarà comunque una distribuzione come quella di Gauss.

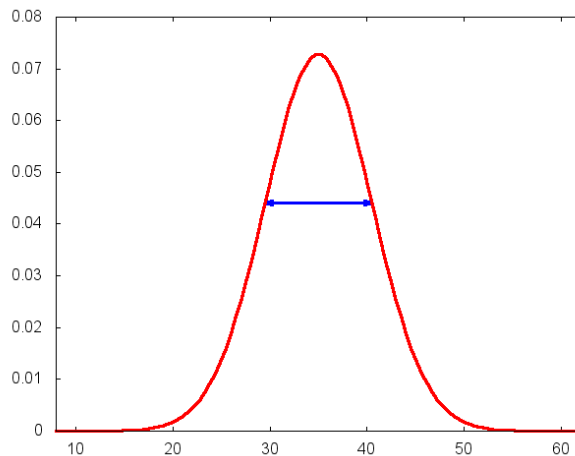


Figura 2.8 Una gaussiana con $\mu = 35$ e $\sigma = 5.48$. La freccia blu indica l'intervallo di $\pm\sigma$ attorno a μ .

Se le misure sono tutte uguali significa che gli effetti casuali che portano a una distribuzione di tipo gaussiano non sono abbastanza grandi da poter essere apprezzati con lo strumento che adoperiamo. Se, per esempio, misurassimo i tempi di oscillazione del mazzo di chiavi con un orologio che mostra al massimo i secondi, probabilmente otterremo sempre il valore 2. Questo significa semplicemente che lo strumento non ha la sensibilità sufficiente per apprezzare gli effetti casuali che possono portare a fluttuare il periodo di oscillazione.

Si può dimostrare che l'area sottesa tra la curva gaussiana e l'asse delle ascisse vale 1 e dal momento che la curva è simmetrica, l'area compresa tra l'asse delle ordinate, quello delle ascisse e la porzione di curva per le $x < \mu$ vale 0.5 ed è uguale a quella sottesa dalla porzione di curva nel semipiano delle $x > \mu$.

2.7 L'interpretazione delle misure

Quando si esegue una misura, questa è sempre affetta da un **errore** o **indeterminazione** o **incertezza** che dir si voglia. Il termine **errore** può essere fuorviante all'inizio perché fa pensare a uno **sbaglio**:

non è così. Una misura è sempre affetta da errore perché è impossibile dire con precisione assoluta quanto valga una grandezza fisica. Le misure infatti si possono eseguire solo attraverso l'uso di qualche strumento che necessariamente deve esprimere il risultato della misura come un numero con un numero finito di cifre. Non esiste alcuno strumento in grado di fornire l'esatta misura dell'area di un cerchio di raggio unitario, che vale π , perché π è un numero che si può scrivere solo con un numero infinito di cifre. Le cifre con le quali possiamo esprimere qualunque misura sono necessariamente finite, perciò non potremo mai sapere, avendo indicato la misura con x , se il suo *vero* valore non sia $x \pm dx$ qualora dx sia più piccolo della sensibilità dello strumento!

Per esempio, se misuriamo il peso⁴ di un pacco di pasta usando una comune bilancia da cucina, possiamo leggere sul display il numero 1.002, se la bilancia può mostrare fino a quattro cifre. Non possiamo sapere se il peso del pacco di pasta è effettivamente 1.0023 o 1.0019, perché la bilancia non può mostrare la cifra in più necessaria per esprimere *correttamente* il peso. Abbiamo quindi un'indeterminazione che potremmo stimare in ± 0.001 unità di misura della bilancia (presumibilmente kg). Il modo corretto di esprimere la misura è dunque

$$M = (1.002 \pm 0.001) \text{ kg} . \quad (2.21)$$

Con questa notazione, di fatto, ammettiamo di non conoscere il peso esatto del pacco, ma di essere certi che il peso sia compreso in un intervallo di ampiezza pari a 1 g in più o in meno rispetto al valore centrale di 1.002 kg; in definitiva stiamo dicendo che sappiamo *soltanto* che il pacco di pasta ha un peso compreso tra $1.002 - 0.001 = 1.001$ kg e $1.002 + 0.001 = 1.003$ kg.

Se misuriamo il peso di più pacchi di pasta, potremmo trovare valori diversi attorno a quello nominale di un chilo. Se osserviamo la forma della distribuzione dei pesi una volta fatto l'istogramma vediamo che i dati si distribuiscono in modo più o meno gaussiano. Il motivo è semplice: nel processo di produzione e confezionamento della pasta è insita

una certa variabilità del tutto casuale su numerosi fattori: la somma di questi effetti casuali porta sempre a una distribuzione gaussiana, qualunque sia la distribuzione delle variabili che provocano le fluttuazioni. Si può dimostrare che, per un numero di misure molto alto, al limite infinito, il **valor medio** $\langle m \rangle$ o semplicemente **media** delle misure, definito come

$$\langle m \rangle = \frac{1}{N} \sum_{i=1}^N m_i = \frac{1}{N} (m_1, m_2, \dots, m_N) \quad (2.22)$$

coincide con la posizione del picco della gaussiana. Si dice che la media è un buon **estimatore** del picco della gaussiana. Nell'espressione della media m_i sono le N misure ottenute.

Seguendo l'esempio sopra potremmo dire che non conosciamo di fatto il **peso** dei pacchi di pasta, ma sappiamo che ogni pacco di pasta ha un peso che varia in un intervallo $m_{\min} - m_{\max}$. Perciò potremmo affermare, sapendo che il valor medio dei pesi misurati sta grosso modo al centro dell'intervallo, che il peso dei pacchi di pasta è

$$M = \left(\langle m \rangle \pm \frac{m_{\max} - m_{\min}}{2} \right) \quad (2.23)$$

prendendo la semiampiezza dell'intervallo dei valori trovati come errore sulla misura del peso m . Quest'ampiezza prende il nome di **semidispersione massima**. Nel caso della misura del peso di un singolo pacco di pasta l'espressione $M = (1.002 \pm 0.001) \text{ kg}$ indicava il fatto che avevamo un'incertezza di 0.001 kg sul valore *vero* della misura, ma che in fondo eravamo **certi** del fatto che, misurando un'altra volta proprio quel peso, avremmo ottenuto lo stesso valore. Questo non è più vero nel caso della misura di un pacco di pasta qualunque: non possiamo escludere che, misurando un altro pacco di pasta il suo peso si trovi all'esterno dell'intervallo di semidispersione trovato. Naturalmente la probabilità che questo accada sarà piccola, ma non nulla! In teoria un qualunque pacco di pasta (che potrebbe appartenere a un lotto diverso da quelli usati per ottenere la distribuzione) potrebbe avere

⁴Vale sempre la nota di pagina 18.

un peso qualunque compreso tra $-\infty$ e $+\infty$, sebbene i valori che distano molto dal valore centrale $\langle m \rangle$ dovrebbero essere assai improbabili.

Quello che possiamo sperare di fare in questi casi è fornire un'indicazione di tipo **statistico** dell'ampiezza della distribuzione: non possiamo cioè dire che **sicuramente** il peso di un altro pacco di pasta sarà compreso nell'intervallo dato, ma che lo sarà con una certa probabilità (alta, ma non uguale a uno).

2.8 Analisi statistica delle misure

Riconsideriamo le misure dei tempi di oscillazione del pendolo fatte da Anna e Bruno (voi fate sempre l'esercizio con le misure fatte da voi). La prima misura eseguita da Anna è $T_1 = 1.13$ s. Significa che Anna è **certa** che il valore *vero* della misura da lei eseguita (non del periodo del pendolo) è compreso tra 1.12 e 1.14 perché dovremmo supporre che Anna ha potuto leggere sul display del suo cronometro il numero 1.13 e quindi ha un'incertezza che al massimo è sulla seconda cifra decimale. La sua misura si esprime come

$$T_1 = (1.13 \pm 0.01) \text{ s}. \quad (2.24)$$

Possiamo fare un istogramma delle misure raggruppandole in intervalli ampi 0.05 s. Nel primo intervallo (anche detto **bin**, in inglese) mettiamo il numero di misure comprese tra $T = 1.00$ e $T = 1.05$, che sono 3; nel secondo quelle comprese tra $T = 1.05$ e $T = 1.1$, che sono 9 e così via. Facciamo attenzione a non contare due volte quelle che dovessero capitare al limite di un bin: nell'esempio che stiamo facendo stiamo considerando il limite destro dell'intervallo incluso e quello sinistro escluso. La Tabella 2.4 riporta il numero di misure che capitano in ciascun intervallo.

Per quanto le misure non siano poi tante, si vede subito (Fig. 2.9) che la distribuzione ha una forma che ricorda quella della curva di Gauss. Il valore medio dei tempi è $\langle T \rangle = 1.12$ (notate che si trova

i	T_i [s]	n_i
1	1.05	3
2	1.10	9
3	1.15	10
4	1.20	2
5	1.25	2

Tavola 2.4 Istogramma delle misure di tempo eseguite da Anna. Nella colonna centrale è riportato solo il valore destro dell'intervallo. È evidente che l'estremo sinistro coincide con l'estremo destro dell'intervallo precedente.

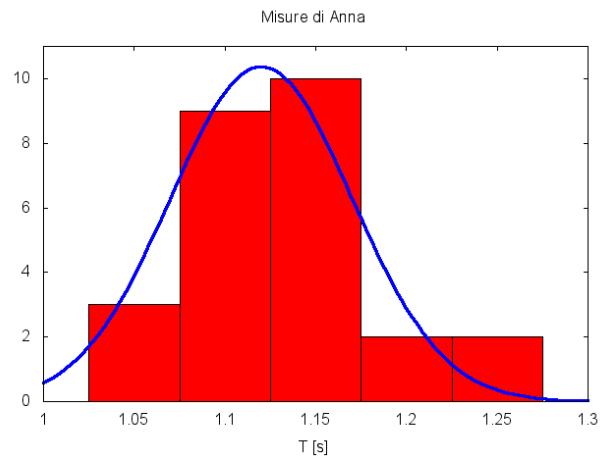


Figura 2.9 Istogramma delle misure eseguite da Anna con la curva di Gauss sovrapposta. La curva è quella descritta nell'Equazione (2.19) moltiplicata per l'area dell'istogramma – l'area della curva di Gauss è pari a 1 – e cioè per 26×0.05).

a metà tra il secondo e il terzo intervallo, dove si trova il picco della distribuzione) e la semidisperzione massima è $\frac{1}{2}(1.25 - 1.02) = 0.12$. Potremmo dunque dire che il periodo del pendolo misurato da Anna è

$$T = (1.12 \pm 0.12) \text{ s}. \quad (2.25)$$

Ma questo vorrebbe dire che siamo **certi** di trovare, eseguendo una nuova misura di tempo, un valore compreso tra 1.00 e 1.24! Basta osservare i dati per capire che non può essere! Almeno una delle misure si trova fuori da questo intervallo: 1.25.

Non possiamo affermare questo, ma possiamo affermare che **molto probabilmente** una nuova misura di tempo fornirà un valore **abbastanza vicino** a 1.12, che è il valor medio. Quanto vicino? E con quale probabilità? Per deciderlo dobbiamo stabilire un criterio per misurare l'ampiezza σ della distribuzione e valutare quanti **eventi** cadono nell'intervallo di ampiezza σ attorno al valor medio.

Una maniera di fornire una misura dell'ampiezza della distribuzione consiste nel misurare quanto, in media, distano le varie misure eseguite dal valore centrale $\langle T \rangle$. La distanza d_i tra la misura i -esima (T_i) e questo valore si esprime come

$$d_i = |T_i - \langle T \rangle| . \quad (2.26)$$

Aver a che fare con l'operatore di **modulo** $|\dots|$ è sempre molto fastidioso, perciò conviene definire la stessa grandezza come

$$d_i = \sqrt{(T_i - \langle T \rangle)^2} \quad (2.27)$$

che sembra più complicata, ma in realtà è più facile da manipolare perché il quadrato è sempre positivo (al contrario della differenza $T_i - \langle T \rangle$) ed estraendone la radice si ottiene lo stesso valore per d_i che si sarebbe ottenuto con il modulo. Il valor medio dei quadrati delle distanze d_i^2 è, per definizione,

$$\langle d^2 \rangle = \frac{1}{N} \sum_{i=1}^N d_i^2 = \frac{1}{N} \sum_{i=1}^N (T_i - \langle T \rangle)^2 \quad (2.28)$$

e la sua radice dovrebbe dare un'idea di quanto sia larga la distribuzione dei pesi m_i . Chiamiamo $\langle d^2 \rangle$ la **varianza** della distribuzione e la indichiamo con σ^2 , così la sua radice σ coinciderà con quella che vogliamo usare per ampiezza dell'intervallo. Chiamiamo σ **deviazione standard**. Per le misure eseguite da Anna la varianza σ^2 vale

$$\sigma^2 = 0.002807 \dots \quad (2.29)$$

Non è molto utile trascrivere tutti i decimali: teniamone solo due, quindi $\sigma^2 \simeq 0.0028$. La deviazione standard quindi vale (sempre tenendo solo due decimali) $\sigma \simeq 0.053$. I dati sono distribuiti, effettivamente, all'interno di un intervallo poco più grande di due volte la deviazione standard. Ma in effetti σ rappresenta bene la **larghezza** della distribuzione: se i dati fossero distribuiti su un intervallo più ampio, anche σ sarebbe più grande. Inoltre, se disegniamo una curva di Gauss con $\mu = \langle T \rangle$ e $\sigma = \sqrt{\langle d^2 \rangle}$, e sovrapponiamo il disegno all'istogramma dei dati, vediamo chiaramente che la deviazione standard $\sigma = \sqrt{\langle d^2 \rangle}$ trovata rappresenta proprio l'ampiezza della curva di Gauss. E del resto, con le tecniche della statistica [2], si può dimostrare che effettivamente $\sigma = \sqrt{\langle d^2 \rangle}$ è un buon estimatore dell'ampiezza della curva di Gauss, nel senso che il suo valore tende sempre di più a quello di una vera gaussiana al crescere di N .

A essere precisi, secondo la statistica, il miglior estimatore della deviazione standard si ottiene partendo dalla varianza calcolata come

$$\langle d^2 \rangle = \frac{1}{N-1} \sum_{i=1}^N (T_i - \langle T \rangle)^2 \quad (2.30)$$

che differisce dal valore calcolato nell'Equazione (2.28) per il denominatore della frazione davanti alla sommatoria che è $N-1$ e non N . La prima (Equazione (2.28)) si chiama **varianza della popolazione**, mentre la seconda (Equazione (2.30)) prende il nome di **varianza del campione**. Per **popolazione** s'intende l'intero insieme dei possibili valori della variabile statistica, mentre per **campione** s'intende una sua parte. Quando eseguiamo delle misure la popolazione è composta di un numero infinito di elementi perché possiamo ripetere la misura quante volte vogliamo. L'insieme delle misure eseguite quindi è un campione dell'intera popolazione di tutti i possibili valori della misura.

In effetti, se N è abbastanza grande la differenza tra la varianza del campione e quella della popolazione è piccola e può essere addirittura trascurabile.

Non ci pare il caso in questa sede di dover entrare nel merito del perché le due definizioni sono diverse per quel -1 , ma possiamo dare una **non dimostrazione** della maggior correttezza della forma (2.30) osservando che per poter stimare una varianza è necessario fare più di una misura! Se $N = 1$ la varianza è di fatto infinita perché non possiamo sapere, ripetendo la misura, quale sarà la probabilità di trovare un valore più o meno vicino a quello ottenuto. L' $N - 1$ a denominatore fa proprio diventare infinita la varianza in questo caso e di fatto **impone** di fare **almeno** due misure.

Se contiamo quanti eventi capitano all'interno dell'intervallo $\pm\sigma$ attorno a $\langle T \rangle$, che va da $1.12 - 0.05 = 1.07$ (stiamo prendendo $\sigma = 0.053 \simeq 0.05$) a $1.12 + 0.05 = 1.17$, troviamo che ce ne sono 17⁵: poco più del 65 % del totale (che è 26). L'area compresa tra la curva di Gauss e l'asse delle ascisse vale 1, e se si calcola quella compresa tra la curva di Gauss e la porzione di asse della ascisse che va da $\mu - \sigma$ a $\mu + \sigma$ si trova che vale circa 0.68. Dal momento che stiamo facendo l'ipotesi che le frequenze approssimino le probabilità, il fatto d'aver trovato 17 eventi su 26 che risiedono all'interno di un intervallo ampio $\pm\sigma$ attorno al valor medio significa che abbiamo una probabilità di circa il 65 % (cioè di circa 2/3) di trovare un valore compreso in questo intervallo rifacendo una misura di tempo.

Se quindi si esprime la misura di tempo come

$$T = (1.12 \pm 0.05) \text{ s} \quad (2.31)$$

s'intende che ci sono 2 probabilità su 3 che, ripetendo una misura di T , si trovi un valore compreso tra 1.07 e 1.17. Prendiamo convenzionalmente questa misura come quella che fornisce a chi legge l'informazione corretta circa la precisione delle misure. Questo significa che dobbiamo aspettarci, in una serie di M misure, che circa 1/3 di queste si trovi fuori dell'intervallo. Si vede facilmente che la probabilità di trovare una misura all'esterno di un intervallo di $\pm 2\sigma$ è del 5 % circa, mentre quella che la misura si

trovi all'esterno di un intervallo ampio 3σ è inferiore all'1 %.

Se questo è il significato che diamo all'**errore** di una misura, allora non possiamo più esprimere il risultato della misura del peso di un pacco di pasta come $M = (1.002 \pm 0.001) \text{ kg}$ perché questo avrebbe un significato diverso (se ripesiamo lo **stesso** pacco di pasta troviamo sempre lo stesso numero e cioè un valore compreso tra 1.001 e 1.003). Perché l'informazione sia la stessa dobbiamo fare in modo che l'errore fornito sulla misura corrisponda a un intervallo che contiene circa 2/3 dei potenziali valori, il che significa che l'errore corretto da assegnare a questa misura è $0.001/3 \simeq 0.0003$ e pertanto la misura si esprime come

$$M = (1.0020 \pm 0.0003) \text{ kg}. \quad (2.32)$$

Avrete notato che abbiamo aggiunto uno zero **non significativo** dopo 1.002 facendolo diventare 1.0020. Questo perché l'errore con il quale conosciamo quel numero si trova sulla quarta cifra decimale, quindi dobbiamo rappresentare il numero in quel modo. In effetti si può dimostrare che la deviazione standard di una variabile casuale con distribuzione uniforme vale $1/\sqrt{12} \simeq 1/3$ l'ampiezza dell'intervallo dei valori permessi.

2.9 Errori sistematici

Gli errori di cui si tratta nel paragrafo precedente sono di natura statistica, ma nell'eseguire una misura si può incappare in un'altra fonte di errore: gli errori **sistematici**. Gli errori statistici provocano fluttuazioni attorno al valore *vero* di una misura che possono essere sia positive che negative, quindi in media sono da considerarsi nulli ed è per questo che assumiamo il valor medio delle misure come il valore *vero* (che scriviamo sempre in corsivo perché il concetto di valore vero non esiste in fisica, dal momento che non possiamo misurarlo).

Gli errori sistematici, invece, sono dovuti a fenomeni che tendono a spostare il valor medio dei risultati rispetto al valore che si avrebbe se tali fonti di errore fossero assenti. Per fare un esempio, se

⁵Dobbiamo sempre scegliere se includere o meno gli eventi ai bordi dell'intervallo, ma per quanto stiamo discutendo la differenza è irrilevante.

si usa un righello le cui tacche siano state incise a una distanza leggermente piú piccola o leggermente piú grande di 1 mm, le misure di lunghezza eseguite con quel righello risulterebbero sempre piú grandi o sempre piú piccole di quelle eseguite con un righello a norma.

Gli errori sistematici possono provenire dalle fonti piú disparate: inaccuratezza degli strumenti (come nel caso del righello), non trascurabilità di effetti spuri (misurando una temperatura si deve tener conto del calore disperso nell'ambiente), influenza dello strumento sulla misura (se lo strumento è freddo, quando si misura la temperatura di un oggetto parte del calore serve a scaldare lo strumento), etc..

Se consideriamo, per esempio, i dati raccolti da Bruno, che ha misurato gli stessi tempi con un cronometro diverso, vediamo che il valor medio delle sue misure è diverso da quello di Anna, essendo $\langle T \rangle = 1.34$, con una deviazione standard di 0.085 (rifate i calcoli). In sostanza le due misure, quella di Anna che indichiamo con T_A e quella di Bruno per la quale usiamo il simbolo T_B sono

$$\begin{aligned} T_A &= (1.12 \pm 0.05) \text{ s}, \\ T_B &= (1.34 \pm 0.09) \text{ s}. \end{aligned} \quad (2.33)$$

Prima di sostenere che i due numeri sono diversi è necessario stabilire che lo siano davvero. In effetti, il fatto che i valori medi siano diversi non ha alcuna importanza: per dire che due numeri sono uguali, in fisica, non si confrontano i valori medi, ma gli intervalli d'incertezza. Infatti, a quanto ne sappiamo, la misura di Anna è un qualunque numero compreso tra 1.07 s e 1.17 s, mentre quella di Bruno è un numero compreso tra $1.34 - 0.09 = 1.25$ s e $1.34 + 0.09 = 1.43$ s. Se gli intervalli si sovrappongono anche parzialmente le due misure sono **compatibili**: rappresentano cioè lo stesso valore. Nel caso in esame il valore piú alto compatibile con le misure di Anna è 1.17 s, mentre quello piú basso trovato da Bruno vale 1.25 s. Le due misure di tempo sono effettivamente tra loro **incompatibili**. Questo naturalmente non è possibile: entrambi hanno misurato gli stessi tempi, quindi le misure possono fluttuare, ma devono essere compatibili!

Il fatto è che uno dei due deve aver commesso qualche **errore sistematico** che sposta tutta la distribuzione sulla destra (se la colpa è di Bruno) o sulla sinistra (se è di Anna). Ci possono essere molte ragioni per cui si verifica un'eventualità del genere e questo esempio chiarisce perché un fisico non si accontenta mai dei risultati ottenuti da un solo gruppo di ricercatori, ma pretende che la stessa misura sia ripetuta da piú persone prima di poterla considerare affidabile.

Il controllo incrociato è indispensabile proprio per valutare l'eventuale presenza di errori sistematici (la presenza di una discrepanza implica l'esistenza di qualche sistematica, ma non è vero il viceversa: se due misure sono compatibili non è affatto detto che non siano affette da qualche tipo di errore sistematico).

Non esistono tecniche standard per la valutazione dell'errore sistematico, proprio perché può derivare dalle cause piú disparate. L'unica cosa che si può fare è cercare di stimarlo e di ridurlo al minimo. Una prima stima molto semplice consiste proprio nel confrontare le misure fatte da sperimentatori e con strumenti diversi.

Non possiamo stabilire chi tra Anna e Bruno abbia introdotto una qualche sorgente di errore sistematico nella misura (forse tutti e due). È chiaro che c'è di mezzo un qualche errore sistematico che, a differenza di quello statistico che può essere solamente ridotto entro certi limiti, si può eliminare attraverso un'accurata progettazione dell'esperimento. In questi casi è necessario eseguire altre misure, controllando tutti i parametri che possono dar luogo all'insorgere di qualche effetto che *spinge* il valore della misura in una direzione piuttosto che nell'altra. Ad esempio, si deve provare a eliminare l'effetto soggettivo di chi acquisisce le misure impiegando un sistema automatico di acquisizione dati (che grazie alle moderne tecnologie è ormai facilissimo realizzare a costo praticamente zero).

L'aver fatto almeno due misure ci permette di valutare l'ordine di grandezza dell'errore sistematico che si può stimare (la stima è necessariamente molto rozza) come grosso modo

$$\sigma_{SYS} \simeq T_B - T_A = 1.34 - 1.12 \simeq 0.2 \text{ s}. \quad (2.34)$$

In ogni caso non abbiamo alcuna ragione per sostenere che una misura è migliore dell'altra (quella di Anna ha l'errore statistico minore, ma questo non vuol dire nulla dal punto di vista dell'accuratezza della misura). Pertanto non possiamo e non dobbiamo **scartare** nessuna delle due misure. Entrambe le misure sono valide se non si identifica in modo chiaro una causa di errore sistematico nella procedura seguita durante l'esecuzione delle operazioni di misura.

Molti studenti alle prime esperienze di laboratorio sono portati a **scartare** misure incompatibili con quanto si aspettano, senza alcuna vera ragione. Questo è un grosso errore. Un buon fisico (in generale un buono scienziato) considera sempre con molta attenzione tutti i risultati sperimentali: anche quelli palesemente in contrasto con la teoria. Prima di decidere che una misura è fatta male bisogna indagare a fondo per essere certi che vada fatto. Se volete essere buoni fisici non dovrete credere a una sola delle parole scritte su questo testo! Intendiamoci: quello che è scritto qui (salvo errori) è vero e corretto fino a quando qualcuno non dimostrerà **sperimentalmente** il contrario! Ma non si deve ciecamente credere nelle leggi fisiche come se fossero dettate da un'Autorità o da una divinità! Le leggi fisiche sono il risultato dell'analisi dei fatti sperimentali e come tali possono cambiare, se cambia il panorama dei dati a disposizione.

Se **Penzias** e **Wilson**, gli scopritori della **radiazione cosmica di fondo** avessero trascurato di approfondire perché l'antenna da loro realizzata producesse un segnale non previsto, bollandolo semplicemente come qualcosa da scartare perché non previsto e non capito, non avrebbero vinto il **Premio Nobel**. Lo stesso vale per molte scoperte fondamentali che si sono potute fare solo perché qualcuno le cui misure non erano perfettamente compatibili con la teoria, invece di attribuire la discrepanza a non meglio precisate fonti d'errore sistematico, indagò a fondo per capire.

2.10 Propagazione degli errori

Se ci pensate bene, il valor medio del periodo trovato da Anna è molto più preciso di quanto sembri a prima vista. Infatti, quando scriviamo che $T = (1.12 \pm 0.05) \text{ s}$ affermiamo che, ripetendo una singola misura troveremmo un valore che, con il 70 % circa di probabilità (2/3) è nell'intervallo di $\pm 0.05 \text{ s}$ attorno a 1.12 s . Ma eseguendo altre 26 misure otterremmo un valor medio che presumibilmente sarà diverso da 1.12 , ma non di molto. Sicuramente non di 0.05 ! È molto improbabile che il valor medio di altre 26 misure sia un numero vicino a 1.17 . È più facile che somigli a 1.12 !

È l'effetto della propagazione degli errori, che discutiamo in questo paragrafo. Eseguire molte misure di una stessa grandezza fisica è certamente utile per valutare l'errore con il quale si conosce quella grandezza fisica, ma è utile anche per ridurre l'incertezza sul valore medio di quella grandezza fisica.

Provate a fare l'esperimento misurando dieci volte una grandezza fisica qualunque (il numero di scarpe di dieci vostri compagni). Calcolate la media e la deviazione standard. Poi scegliete altri dieci compagni e rifate la misura: le singole misure saranno per lo più contenute all'interno dell'intervallo di ampiezza $\pm \sigma$ attorno al valor medio, ma il valor medio di queste dieci misure è molto vicino a quello ottenuto nella prima serie. Se poi riuscite a procurarvi altri 80 numeri di scarpe potete calcolare dieci valori medi prendendo questi numeri a gruppi di dieci. L'istogramma delle dieci misure del valor medio somiglierà a una gaussiana con una larghezza pari a circa 1/3 rispetto a quella delle gaussiane di ogni singolo gruppo. In altre parole, con la notazione $T = (1.12 \pm 0.05) \text{ s}$ comunichiamo ai nostri colleghi qual è la variabilità di T nelle condizioni in cui abbiamo fatto l'esperimento, ma se vogliamo comunicare la precisione con la quale conosciamo il valor medio dobbiamo usare un errore più piccolo. Per capire quanto più piccolo dobbiamo imparare a valutare come si **propaga** un errore su una misura.

Cominciamo con un caso semplice: supponiamo di voler sapere quanto tempo occorre al pendolo di Anna per fare quattro oscillazioni. Chiaramente questo

tempo è

$$T_4 = 4T_1 = 4 \times 1.12 = 4.48 \text{ s.} \quad (2.35)$$

Dal momento che conosciamo T_1 con un'incertezza di 0.05 s, il valore di T_1 potrebbe benissimo essere 1.17 invece che 1.12. In questo caso avremmo

$$T_4 = 4T_1 = 4 \times 1.17 = 4.68 \text{ s.} \quad (2.36)$$

Naturalmente potrebbe anche essere $T_1 = 1.07$, per cui

$$T_4 = 4T_1 = 4 \times 1.12 = 4.28 \text{ s.} \quad (2.37)$$

In altre parole T_4 sarà un numero compreso tra 4.28 s e 4.68 s. L'intervallo entro il quale può variare il suo valore è ampio $4.68 - 4.28 = 0.40$ s: esattamente quattro volte l'ampiezza dell'intervallo d'incertezza di T_1 . In generale possiamo dire che l'errore da attribuire a una grandezza che si calcola come $y = Cx$ dove C è una costante e x il risultato di una misura è

$$\sigma_y = C\sigma_x \quad (2.38)$$

con ovvio significato dei simboli, così T_4 va espresso come

$$T_4 = (4.5 \pm 0.2) \text{ s,} \quad (2.39)$$

essendo $4 \times 0.05 = 0.2$. Notate che abbiamo approssimato il valor medio a un numero che si scrive con una cifra dopo la virgola perché l'errore è su quella cifra decimale.

Quest'osservazione ci porta subito a escogitare una maniera per misurare questo tempo con precisione migliore: se invece di misurare il periodo di un pendolo Anna avesse misurato il tempo necessario a compiere n oscillazioni (per esempio $n = 5$), avrebbe trovato il valore di T_n con un errore σ_n e avrebbe potuto trovare il valore di T_1 dividendo il valore misurato e il suo errore per n . Attenzione a non farvi ingannare: questo non riduce arbitrariamente l'errore perché la misura, per esempio, di cinque oscillazioni, fluttua più della misura di una sola oscillazione perché il tempo durante il quale gli effetti casuali intervengono è quintuplicato; però

misurare T_5 è più facile che misurare T_1 – quanto meno è più *rilassante* – e l'errore, globalmente, diminuisce, anche se non di un fattore 5! In altre parole, se misuriamo T_5 con un errore σ_5 possiamo ricavare T_1 come $T_5/5$ e l'errore come $\sigma_5/5$. Quindi, il semiperiodo del pendolo di Anna vale

$$T_{1/2} = (0.56 \pm 0.03) \text{ s,} \quad (2.40)$$

avendo moltiplicato T e σ per $1/2$ e avendo approssimato il valore di 0.025 a 0.03. Potevamo aspettarcelo, in effetti: se abbiamo una misura con un errore, *stirando* o *contraendo* la misura l'errore si espande o si contrae dello stesso fattore.

Ma tutti sappiamo che moltiplicare un numero x per un numero y equivale a sommare y volte il numero x : così $4T = T + T + T + T$, quindi ci si potrebbe aspettare che l'errore sulla somma di quattro misure di tempo sia il quadruplo dell'errore sulla singola misura. Proviamo. Dividiamo le 26 misure di Anna in 6 gruppi di quattro misure (lasciando fuori due misure). Facendo riferimento alla Tabella 2.2, sommiamo tra loro le due misure di Anna su due righe consecutive (quindi, per esempio, il primo numero si ottiene sommando $1.13 + 1.06 + 1.10 + 1.05$). I valori che otteniamo sono: 4.34, 4.27, 4.48, 4.42, 4.54 e 4.67, la cui media vale $T'_4 = 4.45$. La deviazione standard di ciascuna singola misura, per quanto sopra vale 0.05 quindi ci si potrebbe attendere una deviazione standard attorno a 0.2 per questi valori. Se la calcoliamo otteniamo

$$\sigma_{T'_4} = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (T'_{4i} - \langle T'_4 \rangle)^2} \simeq 0.14 \quad (2.41)$$

che è circa la metà! In effetti, secondo la statistica, ci si aspetta che l'errore sulla somma di questi quattro numeri sia proprio due volte l'errore sul singolo numero e non quattro. Proviamo a capire perché con un esempio nel quale si sommano due numeri x e y . Se x e y sono affetti da errore e con essi si costruisce la grandezza fisica $z = x + y$ potrà accadere che x fluttui fino a $x + \sigma_x$ e che y fluttui fino a $y + \sigma_y$ per cui z sarà $z = x + y + \sigma_x + \sigma_y$ e, in effetti,

fluttua rispetto al valore centrale z di una quantità pari alla somma degli errori su x e y . Ma questo non accade sempre! Molte volte succederà, ad esempio, che x fluttui fino a diventare $x + \sigma_x$, ma y potrebbe fluttuare fino a diventare $y - \sigma_y$ portando z ad assumere il valore $z = x + y + \sigma_x - \sigma_y$ e quindi a fluttuare meno. A differenza del caso in cui la stessa misura è moltiplicata per un numero, in cui l'errore si moltiplica per la stessa quantità, nel caso in cui si sommano due misure gli errori si possono compensare e portare a una riduzione dell'errore statistico complessivo. Per capire come si sommano statisticamente gli errori ricorriamo a una *dimostrazione* di tipo grafico. Se $z_1 = x_1 + y_1$ è il risultato della somma di due misure, z_1 si può rappresentare come un punto di coordinate (x_1, y_1) su un piano cartesiano. Una nuova misura di x e di y condurrà a due valori (x_2, y_2) che individuano un altro punto. Facendo molte misure, ognuna rappresenterà un punto e tutti i punti si distribuiranno attorno al punto medio in modo tale che la maggior parte di queste misure sia compresa all'interno di una regione di piano che dista dal punto medio una stessa quantità: questa regione quindi ha la forma di un cerchio. Il raggio di questo cerchio determina l'errore sulla grandezza fisica z . Avremo dunque che

$$\sigma_z^2 = \sigma_x^2 + \sigma_y^2 \quad (2.42)$$

quindi gli errori statistici si sommano **al quadrato o in quadratura**⁶. Ecco perché sommando quattro numeri con lo stesso errore si ottiene un numero che ha un errore pari a due volte l'errore sulla singola quantità. L'errore è infatti tale per cui la varianza $\sigma_{T_4}^2$ vale

$$\sigma_{T_4}^2 = 4\sigma_T^2 \quad (2.43)$$

per cui

$$\sigma_{T_4} = 2\sigma_T. \quad (2.44)$$

Lo stesso evidentemente vale nel caso della sottrazione tra due grandezze fisiche: basta considerare

⁶Questo comportamento vale solo se la misura di y non è **correlata** con quella di x , cioè se la misura di y non dipende da quella di x .

che le grandezze di cui abbiamo parlato finora x e y possono essere definite negative, quindi anche l'errore su $x - y$, con $x > 0$ e $y > 0$, ha come errore la radice di $\sigma_x^2 + \sigma_y^2$.

A questo punto abbiamo tutti gli ingredienti per valutare l'errore sul valor medio del periodo del pendolo misurato da Anna. Il valor medio di una grandezza fisica si trova sommando N misure e dividendo il risultato per N . L'errore sulla somma delle N misure è pari a $\sqrt{N}\sigma$, assumendo che tutte le misure abbiano lo stesso errore σ . Dividendo questo numero per N l'errore diventa $\sigma/\sqrt{N}$. Di conseguenza, quando si calcola un valor medio, l'errore con il quale lo si conosce è un fattore $\sqrt{N}$ volte più piccolo dell'errore con cui si conosce la singola misura.

Perciò Anna conosce il periodo del pendolo con un errore di $0.05/\sqrt{N} = 0.05\sqrt{26} \simeq 0.01$. Se Anna quindi rimisurasse una volta il periodo del pendolo lo troverebbe per il 70 % delle volte all'interno di un intervallo di ampiezza ± 0.05 attorno al valor medio, ma se misurasse un altro valor medio (cioè se misurasse altre 26 volte il periodo del pendolo e ne traesse il valor medio) questo si discosterebbe dal valore di 1.12 per non più di 0.01 per la maggior parte delle volte. Sarebbe cioè compreso tra 1.11 e 1.13.

Si potrebbe pensare che, in questa maniera, si potrebbe ridurre arbitrariamente l'errore sulla misura di una grandezza fisica! Basterebbe infatti fare in modo che il numero di misure tenda a infinito per portare l'errore a zero: se è vero che non si possono fare infinite misure è vero che se ne possono fare un milione! In realtà, per valutare correttamente l'errore sulla grandezza fisica T del nostro esempio dobbiamo considerare che questo errore ha due sorgenti: una deriva dalle inevitabili fluttuazioni della misura dovute a effetti casuali e l'altra alla sensibilità finita dello strumento con il quale si eseguono le misure, che nell'esempio è di un centesimo di secondo. Quello che possiamo pensare di ridurre mediando i valori è solo il primo: se misurassimo il periodo del pendolo usando un orologio da polso con una sensibilità di 1 minuto troveremmo che il periodo è sempre pari a 1 minuto e ripetere 1 000 volte questo esperimento non aiuta per niente! L'errore che compete alla mi-

sura infatti si valuta sommando (in quadratura) le varie sorgenti di errore e perciò è, più precisamente

$$\sigma_{TOT} = \sqrt{\frac{\sigma^2}{N} + \delta_T^2}, \quad (2.45)$$

avendo indicato con σ la deviazione standard delle misure e con δ_T la sensibilità dello strumento. Visto che nel nostro caso $\delta_T = 0.01 \ll \sigma/\sqrt{N}$, possiamo trascurare questo termine, ma se N diventasse molto grande non potremmo più farlo.

Esercizio 2.2 *La distribuzione della media*

Il fatto che la distribuzione dei valori medi di una variabile casuale abbia una deviazione standard più piccola di quella della singola variabile è un risultato generale, che vale per ogni variabile casuale, comunque distribuita. Fate la prova misurando grandezze qualunque, come la lunghezza delle dita delle mani dei vostri compagni: misurate la lunghezza di ogni dito di entrambe le mani e fate un istogramma delle misure ottenute ricavando, per ogni compagno, la lunghezza media delle dita. Quindi fate un istogramma dei valori medi ottenuti per ciascuno studente e misurate valor medio e deviazione standard di questi. Vedrete che la deviazione standard delle medie è più piccola di un fattore $\sqrt{N}$ rispetto alla deviazione standard di una singola misura (cioè della deviazione standard delle lunghezze delle singole dita).

2.10.1 La media pesata

La **media aritmetica** che si calcola come

$$\langle x \rangle = \frac{1}{N} \sum_{i=1}^N x_i \quad (2.46)$$

è un buon estimatore della posizione del picco della gaussiana solo quando tutte le misure x_i hanno *pari dignità*. Se infatti una o più delle x_i fossero *meno affidabili* di altre, queste dovrebbero influenzare meno il valore della media, che dovrebbe essere determinato in modo che le misure più affidabili **pesino** di più. Una misura è più affidabile dell'altra

(in assenza di errori sistematici), quando il suo errore statistico è minore. Perciò un modo per **pesare** l'importanza di ciascuna misura consiste nel fare in modo che quelle che hanno errore maggiore entrino nella somma moltiplicate per un fattore che ne riduca l'importanza. Se si **pesa** ciascuna misura per un fattore w_i , che dipende dall'errore su x_i , la somma delle misure si scrive come

$$\sum_{i=1}^N w_i x_i. \quad (2.47)$$

Se i pesi fossero tutti uguali (in particolare se $w_i = 1$ per ogni i) il risultato finale dovrebbe coincidere con quello della media aritmetica e quindi dobbiamo imporre che

$$\frac{1}{C} \sum_{i=1}^N w_i x_i = \frac{1}{N} \sum_{i=1}^N x_i. \quad (2.48)$$

Se i pesi sono tutti uguali $w_i = w$ per ogni i e si può portare il simbolo fuori del segno di sommatoria per ottenere

$$\frac{w}{C} \sum_{i=1}^N x_i = \frac{1}{N} \sum_{i=1}^N x_i. \quad (2.49)$$

Quindi $C = Nw = \sum_i w$. Di conseguenza possiamo scrivere che la **media pesata** delle misure x_i è data da

$$\langle x \rangle = \frac{\sum_{i=1}^N w_i x_i}{\sum_{i=1}^N w_i}. \quad (2.50)$$

Notiamo che per $w_i = 1$ la formula sopra si riduce a quella della media aritmetica, come dev'essere. Come peso possiamo prendere qualunque espressione che dipenda dall'errore e che, all'aumentare di questo, faccia diminuire il peso. Per esempio

$$w_i = \frac{1}{\sigma_i^2}. \quad (2.51)$$

Questa scelta permette di semplificare la trattazione che porta alla valutazione dell'errore che si deve attribuire alla media pesata. Nel caso della media aritmetica avevamo che

$$\sigma_{\langle x \rangle} = \frac{1}{N} \sqrt{\sum_{i=1}^N \sigma_i^2} \quad (2.52)$$

che quando tutte le σ_i sono uguali diventa

$$\sigma_{\langle x \rangle} = \frac{1}{N} \sqrt{N \sigma^2} = \frac{\sqrt{N} \sigma}{N} = \frac{\sigma}{\sqrt{N}}. \quad (2.53)$$

Nel caso della media pesata, l'errore da attribuire a ogni addendo della somma a numeratore è

$$w_i \sigma_i = \frac{1}{\sigma_i^2} \sigma_i = \frac{1}{\sigma_i} = \sqrt{w_i}, \quad (2.54)$$

Nella somma ciascuno si somma **in quadratura** con gli altri, per cui la varianza del numeratore si scrive come

$$w_1 + w_2 + \dots \quad (2.55)$$

La deviazione standard del numeratore quindi vale $\sqrt{w_1 + w_2 + \dots}$. Per ottenere la deviazione standard della media pesata dobbiamo dividere questa per la somma dei pesi, che vale $w_1 + w_2 + \dots$ e in definitiva

$$\sigma_{\langle x \rangle} = \sqrt{\frac{1}{\sum_{i=1}^N w_i}}. \quad (2.56)$$

Se, ad esempio, volessimo conoscere la media pesata delle misure di tempo fatte da Anna e da Bruno, che valevano

$$\begin{aligned} T_A &= (1.12 \pm 0.05) \text{ s}, \\ T_B &= (1.34 \pm 0.09) \text{ s}. \end{aligned} \quad (2.57)$$

dobbiamo eseguire le seguenti operazioni: per la media

$$\langle T \rangle = \frac{\frac{1.12}{0.05^2} + \frac{1.34}{0.09^2}}{\frac{1}{0.05^2} + \frac{1}{0.09^2}} \simeq 1.17 \text{ s}, \quad (2.58)$$

mentre per la deviazione standard della media avremmo

$$\sigma_{\langle T \rangle} = \sqrt{\frac{1}{\frac{1}{0.05^2} + \frac{1}{0.09^2}}} \simeq 0.04 \text{ s}. \quad (2.59)$$

Occorre prestare attenzione al significato che si dà a questo valore: 0.04 s non è l'incertezza con la quale conosciamo il periodo del pendolo grazie alle misure di Anna e di Bruno, perché queste misure sono quasi certamente affette da un errore sistematico. Da questo punto di vista questo numero ha un significato meramente matematico: dal momento che i due numeri (diversi) sono conosciuti con un'incertezza rispettivamente di 0.05 s e di 0.09 s, la loro media è nota con un'incertezza di 0.04 s. Ma questa media non ha alcun significato fisico nel caso in esame. Lo avrebbe soltanto se le due misure di tempo fossero tra loro compatibili.

Definire le grandezze fisiche

Una corretta definizione delle grandezze fisiche è importante per riuscire a trarre informazioni dalle misure. Come le leggi fisiche, anche la definizione delle grandezze e delle rispettive unità di misura può cambiare col tempo e con il progredire della conoscenza o delle esigenze. Il tempo, per esempio, è una grandezza fisica che conosciamo tutti senza bisogno di darne una definizione¹, ma gli uomini hanno trovato utile misurarlo per scandire i ritmi della giornata. Inizialmente la misura era estremamente rozza: bastava distinguere il giorno dalla notte. Poi si è presentata l'esigenza di una scansione più precisa delle ore della giornata, per cui il tempo era definito, sostanzialmente in maniera arbitraria, dal sagrestano che suonava la campana della chiesa in occasione delle funzioni religiose che si svolgevano in vari momenti della giornata, definiti da strumenti come le **meridiane** che si basavano sull'ombra di un'asta (**gnomone**) proiettata dal Sole su una parete. La necessità di una misura più accurata portò a riconoscere che il giorno non aveva sempre la stessa durata e a definire un **giorno medio** diviso in 24 ore tutte uguali che si potevano misurare grazie a strumenti semplici come le clessidre o gli orologi ad acqua. L'intensificarsi dei viaggi (sopra tutto quelli per mare) portò allo sviluppo dei primi orologi meccanici, perfezionati con l'avvento del treno, che imponeva una sincronizzazione tra le ore di due stazioni (quella di arrivo e quella di partenza). Solo recentemente si riconobbe la necessità dell'adozione prima di un'orario comune per tutti i paesi e le città di una stessa nazione e poi di una stessa regione del globo terrestre. I **fusi orari** furono adottati nel 1879

in occasione di una conferenza internazionale promossa dai responsabili delle ferrovie di vari Paesi. Molti divertenti aneddoti in proposito si trovano su una **pubblicazione** dell'osservatorio di Arcetri [4]. Di particolare interesse il fatto che il primo a proporre l'adozione dei fusi fu l'italiano **Quirico Filopanti**.

Da allora la definizione di tempo è cambiata molte volte, per adattarsi alle esigenze di natura tecnologica. Oggi una misura di tempo accurata è fondamentale per molte applicazioni: per il funzionamento dei **navigatori satellitari**, ad esempio, è richiesta una precisione dell'ordine dei milionesimi di secondo!

In questo capitolo cominciamo ad analizzare la definizione di alcune grandezze fisiche tra quelle meno ovvie. L'analisi approfondita della loro definizione ci porterà a fare le prime **scoperte**.

3.1 Massa e Peso

Quando diciamo che qualcosa è più **pesante** di un'altra intendiamo dire che, per sollevare la prima si fa più fatica rispetto a quanta se ne fa per sollevare la seconda. Potremmo quindi dire che la misura di peso è una misura della **fatica** che facciamo per sollevare qualcosa. Sfortunatamente questa definizione è troppo vaga e soggettiva: di sicuro l'autore di questa pubblicazione fa più fatica a sollevare un bilanciere da 50 kg di quanta non ne faccia **Arnold Schwarzenegger**², ma non per questo possiamo aspettarci che il bilanciere abbia un peso diverso per il sottoscritto e per un aitante e giovane culturista! Dovremmo farci un'idea più precisa di cosa

¹A questo proposito è interessante leggere il Cap. 13 del Libro XI delle "Confessioni" di **Agostino** [3].

²o forse quanta ne avrebbe fatta qualche anno fa.

intendiamo per **peso** se vogliamo che diventi una grandezza fisica.

Se prendessimo un recipiente e lo riempiamo di una certa sostanza, mantenendo fisso il volume del recipiente, potremmo quanto meno confrontare due pesi: se è vero che la fatica che si fa a sollevare lo stesso peso è soggettiva, tutti sono d'accordo nell'ammettere che una bottiglia da 1 ℓ piena d'acqua pesa meno della stessa bottiglia piena di sabbia o di mercurio.

Una misura è sempre un'operazione di confronto quindi si potrebbe definire il peso di un litro d'acqua come l'unità di misura e poi si potrebbe stabilire quanto volume di un'altra sostanza è necessario per produrre lo stesso effetto di un litro d'acqua. In questo caso la misura di fatica è soggettiva, ma il confronto no. Se la misura di fatica soggettiva ci sembra troppo poco precisa possiamo eseguire il confronto ponendo le due quantità di sostanza sui piatti di una **bilancia** (non quella da cucina o pesapersone, ma un'asta incernierata al centro che deve rimanere in equilibrio se alle sue estremità si pongono due pesi uguali: solo così possiamo fare un vero **confronto**).

Quest'osservazione potrebbe portarci già a una prima definizione coerente del peso come la **quantità di sostanza contenuta in un volume**. Un mattone pieno pesa più di un forato perché, a parità di volume, nel primo c'è più *materia* e il peso dovrebbe essere una misura di questa quantità. Del resto un mattone pieno pesa anche più di un panetto di spugna per fioristi, e questo significa che la materia di cui è fatto il panetto deve avere un peso **intrinsecamente** minore di quello della materia di cui è fatto il mattone. Il peso di qualcosa deve quindi dipendere dal tipo di materiale di cui è fatto.

Ma se diamo questa definizione di peso abbiamo un problema non appena cambiamo le condizioni in cui eseguiamo la misura. Se prendiamo due mattoni uguali sicuramente facciamo la stessa fatica a sollevarli e, ponendoli su una bilancia, la portano a stare in equilibrio. Se però eseguiamo lo stesso esperimento in acqua, la bilancia fornisce lo stesso risultato, ma la fatica che si fa a sollevarli è decisamente mi-

nore. Se poi uno lo immergiamo in acqua e l'altro no, quello che si vede è che anche la bilancia non si trova più in equilibrio! Allora? La quantità di materia presente nel mattone non può cambiare solo per averlo immerso nell'acqua. Dev'esserci qualche altra ragione per cui il suo peso cambia.

Dopo questo semplice esperimento siamo portati a pensare che quel che chiamiamo **peso** non è affatto quel che pensavamo inizialmente e cioè una misura della quantità di materia contenuta in un volume: il peso dev'essere qualcosa che dipende certamente da questa, ma non solo; deve dipendere anche dalle condizioni esterne in cui si fa la misura. In altre parole, il peso è la misura di quanto un oggetto sia attratto verso il basso: in certi casi il peso può persino diventare nullo, anche se la quantità di materia non cambia. Basta guardare un filmato della NASA nel quale si vedono gli astronauti fluttuare senza alcun peso quando sono in orbita per rendersene conto.

Il **peso** dunque non è una proprietà dei corpi, anche se dipende chiaramente da quanta (e quale) materia è presente in essi. Se vogliamo definire quest'ultima quantità è necessaria una procedura che permetta di prescindere dagli effetti esterni come la presenza di acqua o di altri fluidi o di condizioni particolari come il trovarsi a bordo della Stazione Spaziale o sulla Luna. E dal momento che stiamo parlando di una grandezza fisica diversa da quella che sarebbe conveniente chiamare **peso**, dobbiamo darle un altro nome. La chiameremo **massa**. Per misurare una massa si procede in questo modo: si sceglie una massa campione come unità di misura e si confronta il peso di questo con quello dell'oggetto di cui si deve determinare la massa, facendo attenzione al fatto che entrambi gli oggetti si trovino immersi nello stesso fluido e che l'esperimento sia fatto **stando fermi**³. Nella Stazione Spaziale, che si muove rapidamente di moto circolare attorno alla Terra, la misura di massa non si può eseguire (o almeno non si può eseguire come descritto).

³Non è strettamente necessario: basterebbe limitarsi a richiedere che la bilancia si muova di moto rettilineo uniforme, ma per iniziare a definire una procedura operativa va bene così.

Dobbiamo dunque distinguere tra **massa** e **peso**: la prima rappresenta la quantità di materia presente in un volume, la seconda l'intensità con la quale i corpi sono attratti verso il basso. E naturalmente dovremo studiare le proprietà dell'una e dell'altra grandezza fisica perché ogni volta che individuiamo una grandezza dobbiamo capirne le proprietà: è questo il modo in cui si derivano le Leggi fisiche.

Quello che possiamo dire della massa è che apparentemente si **conserva**: dividendo in due un mattone (e raccogliendone tutte le briciole se il taglio non è netto), la massa del mattone e della somma delle sue parti è la stessa. Se mescolando due o più sostanze se ne ottiene un'altra, come accade nelle reazioni chimiche, la somma delle masse dei reagenti è uguale alla massa del prodotto della reazione⁴. Apparentemente non c'è modo di distruggere o di produrre massa dal nulla. Sembra che la massa sia qualcosa che rimanga costante in tutto l'Universo.

3.2 La radioattività

Se tuttavia si fanno esperimenti più precisi si scopre che in effetti non è esattamente così: ci sono casi nei quali la massa effettivamente non si conserva, anche se non sono facili da osservare. Come già accennato sopra, in Chimica, ad esempio, facendo reagire alcune sostanze in determinate proporzioni, si ottengono altre sostanze la cui massa è uguale alla somma delle masse dei reagenti. Il carbonato di calcio (CaCO_3), il comune calcare che si trova nell'acqua, si forma mescolando calcio (Ca), carbonio (C) e ossigeno (O) in proporzioni per cui, al fine di avere solo carbonato di calcio al termine della reazione, occorrono 40.078 g di calcio, 12.0107 g di carbonio e 47.9982 g di ossigeno. Se la legge di conservazione della massa fosse esatta, ci aspetteremmo di trovare

$$40.078 + 12.0107 + 47.9982 = 100.0932 \text{ g} \quad (3.1)$$

⁴Questa Legge prende il nome di **Legge di Lavoisier**, dallo scienziato francese **Antoine Lavoisier** che la formulò basandosi sui dati sperimentali.

e invece se ne trovano 100.0869 (lo 0.006 % in meno). E naturalmente nel recipiente in cui è avvenuta la reazione non c'è alcuna traccia né dei componenti né di altri composti. Una parte della massa è **sparita**. Le violazioni della legge di conservazione della massa sono sempre molto piccole ed è perciò possibile assumerne la validità, a meno di non fare esperimenti molto particolari e precisi. Resta però il fatto che non si tratta di una legge **assoluta**.

C'è un altro caso in cui si può verificare una sia pur piccolissima violazione della legge di conservazione della massa: il **decadimento radioattivo**. Il fenomeno consiste in questo: certe sostanze hanno la capacità di **trasmutare** spontaneamente in sostanze diverse. Il **cobalto** è una di queste: una parte di minerale di cobalto tende a trasformarsi in maniera del tutto spontanea in nichel, il che comporta anche una leggerissima variazione della massa del minerale. Anche il potassio, di cui sono ricche le banane, tende a trasformarsi spontaneamente in calcio o in argon: in linea di principio la massa di una banana quindi diminuisce col tempo (ma questa diminuzione è impercettibile con gli usuali strumenti, quindi se tornate a casa con 990 g di banane avendone pagato 1 kg significa che il vostro fruttivendolo vi ha truffato).

Per capire cosa succede in questi casi è necessario eseguire una o più misure che permettano di stabilire con quali caratteristiche avviene il **decadimento**: solo così sarà possibile interpretare il fenomeno.

Una delle misure che possiamo pensare di fare in questi casi consiste nel determinare quanto minerale di cobalto si è trasformato in nichel in un determinato tempo (in questo caso descriviamo un esperimento *virtuale*: le misure reali si fanno diversamente, ma non è qui il caso di entrare in questo tipo di dettagli). Supponiamo di avere, a un tempo $t = 0$, una quantità $M(0) = M_0$ di cobalto, in particolare di quello usato in radioterapia negli ospedali (non tutto il cobalto è radioattivo: solo certe specie). A un tempo t successivo se ne misura una quantità $M(t)$. Si ripete la misura a tempi diversi, per esempio in giorni diversi, e si osserva che al crescere del tempo la massa di cobalto presente nel minerale si

t (s)	$M(t)$ g
0	10.000
110 500	9.995
190 629	9.993
260 990	9.988
342 350	9.983
427 652	9.980
507 777	9.979
587 669	9.976

Tavola 3.1 Dati relativi alla misura di massa di cobalto in funzione del tempo, per un campione iniziale di 10 g di cobalto per radioterapia.

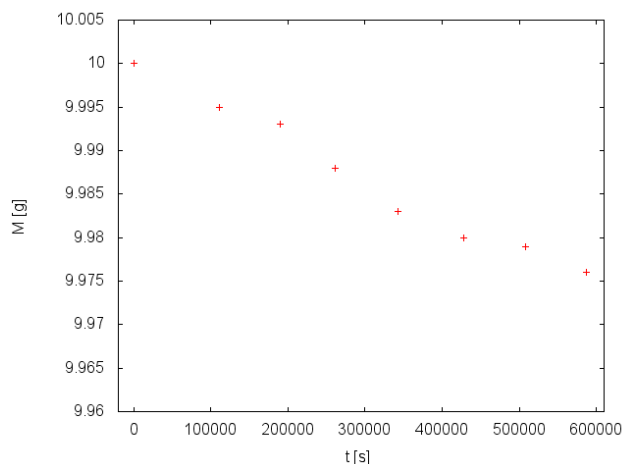


Figura 3.1 La massa di cobalto in funzione del tempo.

riduce sempre di più.

Nella Tabella 3.1 sono riportate alcune misure di questo tipo. Non indichiamo gli errori nella tabella, che assumiamo tutti uguali: questo semplifica la trattazione del problema. Gli stessi dati sono riportati nella Figura 3.1.

Dalla figura si nota una certa tendenza al decremento, sebbene non particolarmente regolare: possiamo ipotizzare che tra la massa del cobalto $M(t)$ e il tempo t sussista una relazione semplice, come quella rettilinea, vale a dire che possiamo pensare di scrivere un'equazione del tipo

$$M(t) = At + B \quad (3.2)$$

dove $A < 0$ rappresenta la pendenza della retta passante per i punti e B l'intercetta. Se si prova a verificare matematicamente che tutti i punti appartengano a una retta si fallisce miseramente. In effetti non possiamo farlo, perché ogni misura è affetta da un errore di cui si deve tenere conto. Sono questi errori che causano le fluttuazioni dei punti sopra e sotto la retta che li descriverebbe in assenza di questi. Come si fa?

3.2.1 La regressione lineare

Non è difficile: basta cercare la retta che si avvicina di più ai dati sperimentali. Dobbiamo però intenderci su che vuol dire **si avvicina**. Una maniera è la seguente: scegliamo un valore di t_i (per esempio $t_4 = 260\,990$ s) e supponiamo che la retta che descrive i dati si possa esprimere come $M(t) = At + B$. Il dato sperimentale corrispondente al valore di t_i scelto, $M_4 = 9.988$ g, può stare un po' sopra o un po' sotto la retta in questione. Quindi dista dalla retta una quantità pari a

$$d_i = |M_i - (At_i + B)|. \quad (3.3)$$

La retta che si avvicina di più a **tutti** i punti è quella che rende minima la distanza

$$d = \sum_{i=1}^N d_i. \quad (3.4)$$

Come nel caso dell'equazione (2.26), aver a che fare con l'operazione di modulo è fastidioso. D'altra parte se d è la minima possibile anche d^2 lo sarà, quindi basterà trovare i valori di A e B che rendono minimo il quadrato di quella distanza. Convenzionalmente questa distanza si indica con la lettera greca χ (pronuncia **chi**), per cui definiamo il χ **quadro** come

$$\chi^2 = \sum_{i=1}^N (M_i - At_i - B)^2. \quad (3.5)$$

Nel caso in esame possiamo assumere che tutti i punti abbiano lo stesso errore pari a 0.001 g, ma se uno dei punti avesse un errore molto più grande degli altri, dovrebbe **pesare** meno nella scelta dei valori di A e B . Per fare in modo che sia così si possono dividere tutti gli addendi per la corrispondente deviazione standard: in questo modo i punti con errore grande contribuiscono poco alla somma. Scriveremo allora che

$$\chi^2 = \sum_{i=1}^N \left(\frac{M_i - At_i - B}{\sigma_i} \right)^2. \quad (3.6)$$

La pendenza e l'intercetta della retta si trovano trovando i valori di A e B che rendono minimo il **chi quadro**. L'operazione che consente di trovare questi valori si chiama **fit** o **regressione lineare**. Espandendo il quadrato si trova

$$\chi^2 = \sum_{i=1}^N \frac{1}{\sigma_i^2} (t_i^2 A^2 + (2Bt_i - 2M_i t_i) A + M_i^2 + B^2 - 2M_i B). \quad (3.7)$$

Allo scopo di semplificare i conti definiamo alcune grandezze:

$$\begin{aligned} S_{tt} &= \sum \frac{t_i^2}{\sigma_i^2} & S_t &= \sum \frac{t_i}{\sigma_i^2} & S_{Mt} &= \sum \frac{M_i t_i}{\sigma_i^2} \\ S_{MM} &= \sum \frac{M_i^2}{\sigma_i^2} & S &= \sum \frac{1}{\sigma_i^2} & S_M &= \sum \frac{M_i}{\sigma_i^2} \end{aligned} \quad (3.8)$$

In questo modo l'equazione (3.7) si riscrive

$$\chi^2 = S_{tt} A^2 + (2BS_t - 2S_{Mt}) A + S_{MM} + SB^2 - 2S_M B. \quad (3.9)$$

che è un polinomio di secondo grado in A , e che quindi rappresenta una parabola il cui vertice ha come ascissa il valore

$$A = \frac{S_{Mt} - BS_t}{S_{tt}}. \quad (3.10)$$

D'altra parte, anche come funzione di B il χ^2 è una parabola:

$$\chi^2 = \sum_{i=1}^N (SB^2 - 2(S_M - AS_t)B + S_{tt}A^2 - 2AS_{Mt} + S_{MM}). \quad (3.11)$$

Il vertice di questa parabola si trova in

$$B = \frac{S_M - AS_t}{S}. \quad (3.12)$$

Per inciso osserviamo che $At + B$ dev'essere una massa quindi le dimensioni di B sono quelle di una massa $[B] = [M]$ e quindi si misura in grammi; anche il prodotto At deve avere le dimensioni di una massa il che significa che A possiede le dimensioni di una massa diviso un tempo e si misura in grammi al secondo (g/s). Corrispondentemente, nell'equazione che fornisce il valore di B per cui il chi quadro è minimo, la differenza $S_M - AS_t$ ha le dimensioni di una massa alla meno uno (S_M è una massa diviso il suo errore al quadrato, quindi è una massa alla meno uno; analogamente per AS_t); diviso per S , che ha le dimensioni di una massa alla meno due, si ottiene una massa⁵. Si verifica facilmente che anche l'espressione di A è corretta dal punto di vista dimensionale.

Dall'ultima equazione scritta ricaviamo A in funzione di B come

$$A = \frac{S_M - BS_t}{S_t} \quad (3.13)$$

che, sostituito nell'equazione che fornisce il valore di A dà

$$\frac{S_M - BS_t}{S_t} = \frac{S_{Mt} - BS_{tt}}{S_{tt}}. \quad (3.14)$$

Risolviamo per B moltiplicando entrambi i membri per $S_t S_{tt}$ per trovare che

$$S_M S_{tt} - B S S_{tt} = S_{Mt} S_t - B S_t S_{tt}. \quad (3.15)$$

e raccogliendo i termini con B a primo membro

⁵I pedici usati nella definizione delle varie somme si possono usare per determinare le dimensioni fisiche delle somme.

$$B(S_t S_t - S S_{tt}) = S_{Mt} S_t - S_M S_{tt} . \quad (3.16)$$

Infine

$$B = \frac{S_{Mt} S_t - S_M S_{tt}}{S_t S_t - S S_{tt}} \quad (3.17)$$

Un rapido controllo dimensionale ci permette di affermare che le dimensioni di B sono quelle attese. Noto il valore di B quello di A si ottiene dall'equazione (3.13). A questo punto non resta che mettere i valori al loro posto. Usando un foglio elettronico, per esempio, è facile costruire la Tabella 3.2. I valori delle celle nelle colonne D ed E e dalla riga 1 alla 8 sono calcolati automaticamente, con una formula. La riga 10 contiene la somma, calcolata automaticamente, dei valori nella colonna corrispondente. A questo punto basta mettere in una cella l'espressione

$$= \frac{E10 * B10 - C10 * D10}{B10 * B10 - 8 * D10} \quad (3.18)$$

dove 8 è il numero di dati che abbiamo a disposizione (poiché tutte le $\sigma_i = \sigma$ nel nostro caso sono uguali, $S = \frac{8}{\sigma^2}$ e i valori di σ^2 si semplificano tra numeratore e denominatore). Si ottiene così il valore

$$B = 10.000 , \quad (3.19)$$

come del resto ci aspettiamo, visto che B deve coincidere (entro gli errori) con $M(0)$. Il valore di A quindi vale

$$A \simeq -42 \times 10^{-9} \text{ gs}^{-1} . \quad (3.20)$$

In sostanza abbiamo stabilito che il nostro minerale di cobalto perde 42 ng ogni secondo (badate: la diminuzione della massa riguarda il cobalto, che si trasforma in nichel, quindi la perdita di massa complessiva del campione è molto più piccola). In un anno ci sono 31 536 000 secondi, quindi la perdita di massa del cobalto è di 1.32 g l'anno. Dal momento che inizialmente avevano 10 g di cobalto, possiamo dire che ogni anno il 13.2 % del cobalto si trasforma in nichel.

La Figura 3.2 mostra i dati con il risultato del fit sovrapposto.

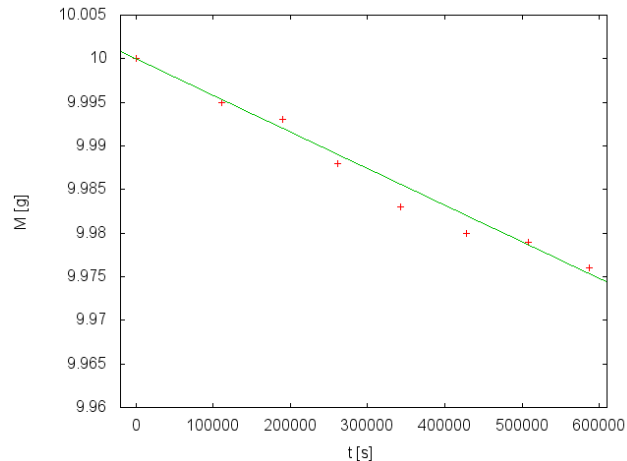


Figura 3.2 I dati relativi al decadimento del cobalto con il risultato del fit sovrapposto.

3.2.2 La costruzione di un modello

Nel modello adoperato per ricavare i parametri del decadimento del cobalto c'è sicuramente qualcosa che non va: se davvero $M(t) = At + B$, a un certo punto succederà che $At = -B$ e quindi $M(t) = 0$ e fin qui tutto bene, ma al crescere di t la massa potrebbe diventare negativa e questo non è ammissibile! È chiaro che la legge del decadimento che abbiamo ipotizzata, seppur verificata con i dati sperimentali, deve essere considerata solo un'approssimazione della vera legge del decadimento.

Proviamo a costruire un **modello** del decadimento radioattivo. Quello che osserviamo sperimentalmente è che certa materia si trasforma spontaneamente in altra materia a un ritmo apparentemente costante, ma che potrebbe non esserlo. Il fatto che la trasmutazione non avvenga istantaneamente per tutto il materiale radioattivo a un certo istante fa pensare che il materiale non si trasforma allo scadere di qualche intervallo di tempo, ma che la trasmutazione procede in maniera progressiva all'infinito.

Facendo misure con diversi campioni si nota anche che la perdita assoluta di massa di materia di un certo tipo aumenta all'aumentare della massa

A	B	C	D	E
i	t_i (s)	M_i (g)	t^2 (s ²)	$M_i t_i$ (gs)
1	0	10.000	0	0.0
2	110 500	9.995	12 210 250 000	1 104 447.5
3	190 629	9.993	36 339 415 641	1 904 955.6
4	260 990	9.988	68 115 780 100	2 606 768.1
5	342 350	9.983	117 203 522 500	3 417 680.1
6	427 652	9.980	182 886 233 104	4 267 967.0
7	507 777	9.979	257 837 481 729	5 067 106.7
8	587 669	9.976	345 354 853 561	5 862 585.9
9				
10	2 427 567	79.894	1 019 947 536 635	24 231 510.9

Tavola 3.2 Una tabella con i dati raccolti, i quadrati dei tempi e i prodotti $M_i t_i$. Alla riga 10 si trovano le somme dei dati nelle rispettive colonne.

iniziale, il che significa che la percentuale di materia che subisce la trasformazione per unità di tempo dev'essere una frazione della massa iniziale, cioè

$$\Delta M \propto -M \quad (3.21)$$

avendo indicato con ΔM la differenza di massa tra l'istante t e l'istante $t = 0$ (la massa diminuisce, quindi questa differenza è negativa: per questo c'è il segno meno a secondo membro dell'equazione). Inoltre ΔM aumenta col tempo apparentemente linearmente, perciò dev'essere anche

$$\Delta M \propto t. \quad (3.22)$$

Dovendo essere ΔM proporzionale sia a M che a t possiamo scrivere che

$$\Delta M = -\alpha M t \quad (3.23)$$

dove α è una costante che dipende solamente dal tipo di materiale radioattivo. Possiamo riscrivere l'equazione come

$$\frac{\Delta M}{M} = -\alpha t. \quad (3.24)$$

All'[appendice matematica](#) si dimostra come questo significhi che M è funzione del tempo e si possa scrivere come

$$M(t) = M(0) \exp(-\alpha t) \quad (3.25)$$

che non è in contraddizione con il nostro modello semplificato, perché per t sufficientemente piccoli, $M(t)$ si può approssimare piuttosto bene con un andamento rettilineo. Per capire quanto piccolo dev'essere t perché l'approssimazione sia valida conviene definire α , che ha le dimensioni fisiche di un tempo alla meno uno, come

$$\alpha = \frac{1}{\tau} \quad (3.26)$$

con τ avente le dimensioni di un tempo e che chiameremo **tempo di decadimento**. In questo modo $M(t)$ si riscrive

$$M(t) = M(0) \exp\left(-\frac{t}{\tau}\right) \simeq M(0) \left(1 - \frac{t}{\tau}\right). \quad (3.27)$$

A questo punto è immediato identificare il parametro A della retta che abbiamo ricavato al paragrafo precedente con

$$A = -\frac{M(0)}{\tau} \quad (3.28)$$

per cui

$$\tau = -\frac{M(0)}{A} = \frac{10}{42 \times 10^{-9}} \simeq 0.23 \times 10^9 \text{ s}. \quad (3.29)$$

Questo tempo corrisponde a circa sette anni e mezzo. Diremo pertanto che il tempo di decadimento del cobalto è di 7.5 anni. Un altro modo di esprimere questo risultato consiste nel valutare quanto tempo è necessario affinché la materia decaduta sia pari alla metà di quella originale, cioè il valore di $t = t_2$ per cui

$$M(t_2) = M(0) \exp\left(-\frac{t_2}{\tau}\right) = \frac{M(0)}{2}. \quad (3.30)$$

Prendendo il logaritmo di entrambi i membri e ricordando che $\log(1/2) = -\log 2$ abbiamo che

$$-\frac{t_2}{\tau} = -\log 2 \quad (3.31)$$

quindi

$$t_2 = \tau \log 2 \simeq 5.2 \text{ anni}. \quad (3.32)$$

Il tempo t_2 si chiama **tempo di dimezzamento** e ci dice quanto dobbiamo attendere prima che il materiale sia decaduto per metà del suo peso iniziale.

Il modello a questo punto va confermato, eseguendo misure a tempi lunghi (rispetto al tempo di decadimento, quindi per periodi di almeno 2–3 anni, fino anche a 10). È anche necessario verificare se il modello si adatta anche ad altre specie radioattive, misurando i tempi di decadimento di altri materiali.

Gli esperimenti dicono che in effetti è così: il modello funziona per tutti i materiali radioattivi e che ognuno di essi ha un suo specifico tempo di decadimento che dipende solo dalla specie (non dipende da altre variabili come quantità iniziale di materiale, luogo di reperimento dello stesso, temperatura, etc.).

In definitiva, solo grazie alla necessità di definire in maniera opportuna una grandezza fisica abbiamo cominciato a fare le prime scoperte! Abbiamo già un primo modello che possiamo considerare come una **Legge fisica**.

3.3 La Temperatura

Al paragrafo precedente abbiamo accennato alla **temperatura**. Anche questa è una grandezza fisica per la quale occorre definire una procedura operativa per la sua determinazione. Cominciamo, come al solito, con qualche semplice esperimento per cercare di definire bene cosa intendiamo per temperatura. La prima cosa che viene in mente è che la temperatura è associata alla sensazione di **caldo** e di **freddo**. È chiaro che non possiamo definire la temperatura in base a sensazioni soggettive: in una classe c'è chi ha sempre freddo, anche in primavera inoltrata e chi ha caldo anche in pieno inverno. Su una cosa però possiamo essere tutti d'accordo: mettendo le mani in due secchi pieni d'acqua possiamo stabilire facilmente quale dei due è caldo e quale è freddo (purché la differenza di temperatura sia abbastanza grande, ma questo è un problema di sensibilità del nostro **strumento** costituito dalle mani). Il secchio di acqua fredda può apparire tale o meno alle diverse persone, ma appare sempre più freddo dell'altro a tutti. Quindi i nostri sensi ci permettono di stabilire se due sostanze hanno la stessa temperatura e, nel caso non ce l'abbiano, quale dei due è più freddo.

Se si mescola l'acqua nei due secchi succede che l'acqua assume una temperatura intermedia tra quelle di partenza: quella calda si raffredda e quella fredda si riscalda. Qualcosa di simile accade se mettiamo in contatto due corpi solidi o mescoliamo due gas.

Possiamo così sostenere che due corpi, posti a contatto tra loro, prima o poi raggiungono una temperatura comune che chiamiamo **temperatura d'equilibrio**. A quest'osservazione si dà talvolta il nome di **principio zero della termodinamica**.

Ora facciamo un altro esperimento: cerchiamo una superficie in legno (la cattedra) e una in metallo (le gambe dei banchi) e tocchiamole. Quale delle due è più calda? La risposta è quasi sempre scontata: **quella di legno**. Basta pensarci un po' per rendersi conto che non è possibile! O meglio, che non è possibile definire la temperatura nella maniera in cui stiamo procedendo. Se infatti è vero il principio zero della termodinamica (e appare essere vero dal momento che tutti gli esperimenti di questo tipo danno il risultato atteso), le gambe e la superficie dei banchi o della cattedra devono trovarsi alla stessa temperatura perché sicuramente sono in contatto tra loro da molto tempo! E allora? Evidentemente il senso del tatto non fornisce una misura della temperatura degli oggetti che tocchiamo: probabilmente misura qualcos'altro.

Facendo esperimenti con corpi caldi e freddi ci si può imbattere in un fenomeno abbastanza curioso: in quasi tutti i casi, un corpo riscaldato aumenta di volume. Poco, ma lo fa. Se non vediamo alcun aumento in certi casi è perché evidentemente non abbiamo la sensibilità necessaria. Se vogliamo capirci qualcosa dobbiamo fare qualche misura e per farle in maniera ordinata dobbiamo usare sistemi che siano semplici, il cui comportamento cioè possa dipendere solo da pochi parametri.

Visto che quel che succede ai corpi riscaldati è che si espandono e che quindi cambiano forma, prendere un oggetto dalla forma complicata non è il caso, evidentemente: conviene lavorare con oggetti la cui forma sia la più semplice possibile! La cosa migliore è partire con qualcosa di molto lungo e stretto, in modo tale che si possa assimilare a una linea priva di spessore. Se quest'oggetto si dilata lo farà prevalentemente lungo la sua lunghezza; le dimensioni trasversali cambieranno in modo impercettibile. Prendiamo quindi un'asta di qualche materiale e immergiamola in acqua fredda. Ne misuriamo la lunghezza, poi la immergiamo in acqua calda e misuriamo nuovamente la sua lunghezza. Se l'asta è di metallo (ferro, rame, alluminio, etc.), quello che si vede dagli esperimenti è che l'allungamento $\Delta\ell = \ell - \ell_0$, dove ℓ è la lunghezza dell'asta calda e ℓ_0 quella dell'asta fredda, è proporzionale alla lunghezza iniziale

dell'asta ℓ_0 . In altre parole più è lunga l'asta immersa in acqua fredda, più si allunga quando si mette nell'acqua calda:

$$\Delta\ell \propto \ell_0. \quad (3.33)$$

Dalle misure di lunghezza si vede anche che usando aste di metalli diversi l'ampiezza dell'allungamento è diversa, quindi la proprietà di allungarsi col caldo dipende dal tipo di materiale. Per questa ragione possiamo ritenere che quando non osserviamo alcun allungamento, come nel caso di aste di materiali come plastica o legno, è perché evidentemente l'entità di quest'ultimo dev'essere molto piccola.

Per quanto i nostri sensi non ci permettano di avere una misura assoluta di temperatura, abbiamo già osservato che sono capaci di trasmettere sensazioni che ci permettono di stabilire in maniera univoca quale, tra due oggetti, sia più o meno freddo dell'altro. In questo modo possiamo disporre di una serie di campioni di acqua a temperature diverse e crescenti e osservare che l'allungamento delle aste di lunghezza iniziale ℓ_0 quando immerse nell'acqua più fredda, si allungano sempre di più man mano che l'acqua diventa più calda. In particolare si vede che l'allungamento percentuale $\Delta\ell/\ell_0$ è lo stesso a parità di differenza di temperatura e a parità di materiale. Possiamo allora **definire** la differenza di temperatura ΔT come qualcosa che provoca l'allungamento delle aste e scrivere che

$$\Delta T = \mu \frac{\Delta\ell}{\ell_0} \quad (3.34)$$

dove μ è un parametro che dipende dal materiale di cui è fatta l'asta. In questo modo due aste uguali si dilatano nello stesso modo se fatte passare da una temperatura all'altra. Se la stessa asta si allunga di più significa che la differenza di temperatura è maggiore. Visto che le aste si allungano passando dal freddo al caldo, $\Delta\ell = \ell - \ell_0 > 0$ di conseguenza anche ΔT dev'essere positiva se attribuiamo a μ questo stesso segno e quindi un oggetto freddo deve avere una temperatura più bassa di uno caldo. L'allungamento subito da un'asta passando da una temperatura a un'altra maggiore è

$$\frac{\Delta \ell}{\ell_0} = \frac{\Delta T}{\mu}. \quad (3.35)$$

Se a parità di differenza di temperatura un'asta si allunga più d'un'altra significa che il parametro μ dell'asta che si allunga di più è più piccolo. Per dare maggiore naturalezza a questo parametro possiamo definire

$$\alpha = \frac{1}{\mu} \quad (3.36)$$

così

$$\frac{\Delta \ell}{\ell_0} = \alpha \Delta T. \quad (3.37)$$

Il parametro α , che come μ dipende dal materiale di cui è fatta l'asta, lo chiamiamo **coefficiente di dilatazione termica**. Maggiore è il valore di α più grande sarà l'allungamento a parità di tutte le altre condizioni.

Se definiamo così la temperatura possiamo eseguire una misura di temperatura facendo, di fatto, una misura di lunghezza! Basta trovare due punti di riferimento cui attribuire un valore arbitrario e il gioco è fatto. Nel corso del tempo sono state fatte le scelte più fantasiose: ad esempio **Gabriel Fahrenheit** definì lo zero della scala termica come la temperatura più bassa che riusciva a raggiungere nel suo laboratorio, mentre come temperatura pari a 100 gradi quella del sangue del suo cavallo (!). Evidentemente una tale scala non poteva sopravvivere all'uso in fisica e fu rapidamente soppiantata dalla scala proposta da **Anders Celsius** che porta il suo nome⁶.

La scala Celsius si basa sulle proprietà dell'acqua. Anche se l'acqua si comporta in modo anomalo rispetto alle altre sostanze: quando è molto fredda, ad esempio, se la si raffredda ulteriormente non diminuisce il proprio volume, ma lo aumenta come è facile constatare mettendo una bottiglia piena d'acqua in un freezer. Nonostante ciò si presta bene al nostro scopo. Il fatto che l'acqua ghiacci ci permette di usare il punto di solidificazione dell'acqua come un

riferimento preciso per attribuire un valore alla sua temperatura. Definiamo la temperatura dell'acqua che sta ghiacciando (o del ghiaccio che sta fondendo) come quella che corrisponde al valore $T = 0$ in un'unità che chiamiamo **grado Celsius**. Indichiamo quest'unità con il simbolo $^{\circ}\text{C}$. Facendo salire la temperatura dell'acqua, scaldandola sul fuoco, a un certo punto questa inizia a bollire. Quando si giunge a questo punto abbiamo un altro chiaro riferimento per attribuire un valore a questa temperatura che definiamo essere quella che corrisponde al valore di $T = 100^{\circ}\text{C}$. Dividendo quest'intervallo in cento parti, l'ampiezza di ciascuno corrisponde a un grado.

Se abbiamo un'asta di alluminio che immersa nel ghiaccio è lunga 1 m, immergendola in acqua bollente si allunga di circa 2.2 mm. Il coefficiente di dilatazione termica dell'alluminio quindi vale

$$\alpha = \frac{\Delta \ell}{\ell_0} \frac{1}{\Delta T} = \frac{2.2 \times 10^{-3}}{1} \frac{1}{100} = 2.2 \times 10^{-5} \text{ } ^{\circ}\text{C}^{-1} \quad (3.38)$$

Di conseguenza un allungamento di 0.022 mm dell'asta corrisponde a un incremento di temperatura di un grado. Così ora abbiamo la possibilità di misurare in maniera oggettiva la temperatura di qualunque cosa: basta metterla in contatto con la nostra asta di alluminio che così diventa un **termometro**, cioè uno strumento per misurare la temperatura. Certo, non è molto pratico, ma almeno il principio lo è e possiamo usarlo per costruire un oggetto più semplice da usare.

Se invece di un'asta rigida usassimo un liquido, anche questo si espanderebbe con la temperatura, ma dovrebbe essere contenuto in un recipiente, che tuttavia potremmo costruire in un materiale con coefficiente di dilatazione molto più piccolo in modo da renderne trascurabile la variazione di volume. Se come recipiente usiamo un tubo lungo e stretto, per quanto esistano liquidi ad alta espandibilità come l'etanolo (il cui coefficiente è $\alpha = 25 \times 10^{-5} \text{ } ^{\circ}\text{C}^{-1}$), possiamo sperare di ottenere espansioni più ampie, ma comunque molto piccole. Possiamo però usare un **trucco**: il liquido lo mettiamo in un recipiente di forma qualunque (per esempio una sfera) di di-

⁶Nell'uso comune però i Paesi anglosassoni hanno conservato, per campanilismo, questa scala, come del resto hanno fatto sulle scale di lunghezza.

mensioni relativamente grandi, collegato a un sottile capillare. Espandendosi, il liquido comincia a risalire lungo il capillare. Ma di quanto? A questo punto bisogna capire quanto aumenta il volume di una sostanza quando è riscaldata. Per capirlo immaginiamo di disporre di dodici sottilissime aste di materiale che si espande con la temperatura, con le quali costruiamo un cubo di cui ogni asta costituisce uno spigolo.

Aumentando la temperatura ogni asta aumenta la propria lunghezza di $\Delta\ell = \ell_0\alpha\Delta T$. Di conseguenza aumenta il volume del cubo. Cerchiamo di calcolare la differenza di volume tra un cubo a temperatura T e uno a temperatura $T' = T + \Delta T$. A temperatura T il volume del cubo è

$$V = \ell_0^3. \quad (3.39)$$

Alla temperatura T' il volume del cubo è

$$V' = \ell^3 = (\ell_0 + \Delta\ell)^3. \quad (3.40)$$

Espandiamo il cubo:

$$V' = \ell_0^3 + \Delta\ell^3 + 3\ell_0^2\Delta\ell + 3\ell_0\Delta\ell^2. \quad (3.41)$$

La differenza $V' - V$ quindi vale

$$V' - V = \Delta\ell^3 + 3\ell_0^2\Delta\ell + 3\ell_0\Delta\ell^2. \quad (3.42)$$

Ora osserviamo che $\Delta\ell = \alpha\ell_0\Delta T$ è un numero piccolo (sicuramente minore di uno) che, se elevato al quadrato, diventa ancora più piccolo, per non parlare di quando è elevato al cubo! Quindi $\Delta\ell^3$ è trascurabile rispetto a $\Delta\ell$. Lo stesso vale per l'addendo proporzionale a $\Delta\ell^2$. Di conseguenza possiamo scrivere

$$V' - V \simeq 3\ell_0^2\Delta\ell \quad (3.43)$$

che divisa per V dà

$$\frac{V' - V}{V} = \frac{\Delta V}{V} = \frac{3\ell_0^2\Delta\ell}{\ell_0^3} = 3\frac{\Delta\ell}{\ell_0} = 3\alpha\Delta T. \quad (3.44)$$

Noto il coefficiente di espansione lineare, quello di volume si ricava moltiplicandolo per tre. Possiamo verificare sperimentalmente questa *legge* che vale per tutti i materiali **isotropi** (che si espandono nello stesso modo in tutte le direzioni). A questo punto abbiamo tutti gli ingredienti per predire cosa succede all'etanolo contenuto in una sfera di raggio R . Il suo volume a una data temperatura alla quale è interamente contenuto nella sfera è

$$V = \frac{4}{3}\pi R^3. \quad (3.45)$$

Se la temperatura aumenta il suo volume aumenta di

$$\Delta V = 3V\alpha\Delta T = 4\pi R^3\alpha\Delta T, \quad (3.46)$$

Se l'espansione può avvenire in un cilindro di raggio r e altezza h dev'essere

$$3V\alpha\Delta T = 4\pi R^3\alpha\Delta T = \pi r^2 h. \quad (3.47)$$

Quindi l'altezza della colonna di etanolo nel capillare vale

$$h = 4\frac{R^3}{r^2}\alpha\Delta T. \quad (3.48)$$

Se $R = 1$ mm e $r = 0.1$ mm l'altezza della colonna è

$$h = 4\frac{1}{0.01}\alpha\Delta T \quad (3.49)$$

e per un aumento di un grado di temperatura l'etanolo sale di

$$h = 4 \times 100 \times 25 \times 10^{-5} \times 1 \text{ m} = 0.1 \text{ m} \quad (3.50)$$

corrispondenti a ben 10 cm!

Ora possediamo uno strumento capace di misurare la temperatura con una sensibilità notevole! Di certo potremo apprezzare differenze di altezza dell'ordine del millimetro, almeno, corrispondenti a un decimo di grado.

Calore e temperatura

Con un **termometro** possiamo misurare le temperature in modo oggettivo. Resta da capire perché percepiamo una sensazione diversa quando tocchiamo il metallo o il legno e cosa provoca le variazioni di temperatura.

Il progresso della Fisica (e delle altre scienze) di solito si snoda attraverso strade piuttosto tortuose che includono errori e interpretazioni sbagliate dei dati sperimentali, che costringono a ripensare i modelli e talvolta a tornare indietro per intraprendere strade che erano state abbandonate. A volte è utile ripercorrere, almeno parzialmente, tali strade, perché **sbagliando s'impara** ed è anche attraverso la conoscenza degli errori fatti dai nostri predecessori che s'impara qualcosa di più non tanto della Fisica, quanto del **Metodo**. È quello che facciamo in questo capitolo.

4.1 La teoria del calorico

Quando mettiamo in contatto due oggetti a temperatura diversa, entrambi raggiungono una temperatura che definiamo **di equilibrio** intermedia tra quelle iniziali dei due corpi. Per capire cosa avviene in questi casi si deve fare un esperimento nel quale si mettono in contatto due sostanze a temperatura diversa. Inizialmente conviene mettersi in una condizione semplice che consiste nel misurare la temperatura di equilibrio raggiunta da due corpi identici, tranne che per la temperatura. Se, ad esempio, prendiamo un bicchiere d'acqua a $T_1 = 20^\circ\text{C}$ e uno a $T_2 = 40^\circ\text{C}$, mescolando insieme i due campioni otteniamo acqua a temperatura $T_{eq} = 30^\circ\text{C}$.

Se cambiamo la temperatura di uno dei due bicchieri la temperatura di equilibrio che si raggiunge

è sempre la temperatura media delle temperature iniziali: $T_{eq} = (T_1 + T_2)/2$.

Proviamo a cambiare la quantità di acqua che mescoliamo con l'altra: per esempio, mescoliamo un bicchiere d'acqua a temperatura T_1 , che contiene una massa d'acqua pari a m_1 , con mezzo bicchiere di acqua a temperatura T_2 di massa pari a m_2 . Se si lascia invariata m_1 e si cambia m_2 tra una misura e l'altra si trova che all'aumentare di m_2 la temperatura di equilibrio è sempre più vicina a T_2 . Se $m_2 = 0$ evidentemente $T_{eq} = T_1$ quindi la legge che determina T_{eq} dev'essere del tipo

$$T_{eq} = \frac{m_1 T_1 + m_2 T_2}{m_1 + m_2}. \quad (4.1)$$

Se $m_2 = 0$, infatti, $T_{eq} = T_1$ e quando $m_2 \gg m_1$ possiamo scrivere che

$$T_{eq} \simeq \frac{m_2 T_2}{m_2} \simeq T_2. \quad (4.2)$$

Basta fare una serie di esperimenti con diversi valori di m_2 per verificare che in effetti è così. Se ora al posto dell'acqua, nel secondo bicchiere mettiamo altre sostanze, vediamo che vale la stessa legge, ma con una differenza: è come se la massa m_2 fosse diversa da quella effettivamente usata. Possiamo interpretare questo fatto assumendo che in effetti la temperatura di equilibrio che si ottiene dipende non solo dalla massa della sostanza, ma anche dal tipo di sostanza.

Se mescoliamo a $m_1 = 0.1$ kg di acqua a $T = 20^\circ\text{C}$ con $m_2 = 0.2$ kg d'acqua a $T = 30^\circ\text{C}$, la temperatura d'equilibrio è

$$T_{eq} = \frac{m_1 T_1 + m_2 T_2}{m_1 + m_2} = \frac{0.1 \times 20 + 0.2 \times 30}{0.1 + 0.2} \simeq 27^\circ\text{C}. \quad (4.3)$$

Se invece di aggiungiamo all'acqua a temperatura T_1 , la stessa quantità d'etanolo a temperatura T_2 la temperatura che la miscela raggiunge è di $T_{eq} \simeq 25^\circ\text{C}$. È come se avessimo introdotto nella miscela meno liquido. Dal punto di vista puramente matematico possiamo scrivere che la temperatura di equilibrio raggiunta dalla miscela si può esprimere come

$$T_{eq} = \frac{m_1 T_1 + c_2 m_2 T_2}{m_1 + c_2 m_2}, \quad (4.4)$$

dove c_2 è un coefficiente (in questo caso minore di uno) che dipende dal tipo di materiale usato nella miscela. Per ogni sostanza possiamo misurare un valore di c_2 che dipende solo dal tipo di sostanza impiegata e possiamo perciò scrivere, in generale, che mescolando due sostanze qualunque i cui coefficienti sono rispettivamente c_1 e c_2 , la temperatura di equilibrio raggiunta sarà pari a

$$T_{eq} = \frac{c_1 m_1 T_1 + c_2 m_2 T_2}{c_1 m_1 + c_2 m_2}. \quad (4.5)$$

Per l'acqua il coefficiente $c_{acqua} = 1$. Cerchiamo ora di dare un senso fisico alle misure che abbiamo fatto. Se mescolando una sostanza calda a una fredda la prima si raffredda e la seconda si riscalda, fino a quando non raggiungono la stessa temperatura, significa che molto probabilmente la sostanza calda **perde** qualcosa che finisce nella sostanza fredda. La perdita di questo qualcosa si manifesta come un abbassamento della temperatura, mentre l'acquisto come un innalzamento. Il passaggio di questo qualcosa dalla sostanza calda alla sostanza fredda si arresta quando si raggiunge un equilibrio nelle temperature.

Tutto fa pensare a un analogo idraulico: se mettiamo in contatto due recipienti contenenti un liquido attraverso un tubo, questo comincia a passare dal recipiente nel quale il livello è più alto a quello nel quale il livello del liquido è più basso. Notate che

non è la quantità di liquido che determina il verso nel quale avviene il trasferimento, ma l'altezza: se un recipiente è lungo e stretto e l'altro largo e basso, nel primo può starci meno liquido che nel secondo, ma se è più alto il liquido si trasferisce da questo all'altro e non viceversa. Il passaggio di liquido si arresta quando il livello nei due recipienti è lo stesso.

Sembrerebbe del tutto logico, a questo punto, fare un'analogia secondo la quale il ruolo dell'altezza del liquido è giocato dalla temperatura, mentre la quantità di liquido contenuto in ciascun recipiente è analogo alla quantità di **qualcosa** che si deve trovare nei corpi in proporzione alla loro massa e secondo il tipo di sostanza. In altre parole sembrerebbe ragionevole pensare che nei corpi sia contenuto una specie di **fluido**, che possiamo chiamare **fluido calorico**, che passa da un corpo caldo a uno freddo fino a quando non s'instaura un equilibrio.

Possiamo individuare la natura di questo fluido analizzando la relazione che ci dà la temperatura d'equilibrio che, moltiplicata per $c_1 m_1 + c_2 m_2$ si riscrive come

$$c_1 m_1 T_{eq} + c_2 m_2 T_{eq} = c_1 m_1 T_1 + c_2 m_2 T_2 \quad (4.6)$$

che si può a sua volta riscrivere portando i termini con m_2 a primo membro e quelli con m_1 a secondo membro:

$$c_2 m_2 (T_{eq} - T_2) = c_1 m_1 (T_1 - T_{eq}). \quad (4.7)$$

Chiamando $\Delta T_1 = T_{eq} - T_1$ la variazione di temperatura della sostanza 1 (cioè la differenza tra la temperatura finale e quella iniziale) e $\Delta T_2 = T_{eq} - T_2$ quella della sostanza 2, l'equazione sopra scritta diventa

$$c_2 m_2 \Delta T_2 = -c_1 m_1 \Delta T_1, \quad (4.8)$$

o, il che è lo stesso,

$$c_2 m_2 \Delta T_2 + c_1 m_1 \Delta T_1 = 0. \quad (4.9)$$

Ciascun addendo rappresenta la variazione di una quantità che si scrive come il prodotto $c_i m_i \Delta T_i$. c_i e

m_i non cambiano nel corso del processo di mescolamento, ma ΔT_i sí e può essere positiva o negativa, secondo che la temperatura finale sia piú alta o piú bassa di quella di partenza. La variazione complessiva di questa quantità è nulla il che si traduce nel fatto che questa quantità è **conservata**, cioè resta costante. La teoria del **fluido calorico** sembra confermata: se pensiamo che la sostanza i possieda una certa quantità iniziale di fluido $Q_i = c_i m_i T_i$, mescolandola con un'altra, quella che ne ha di piú la cede all'altra fino a quando le quantità di fluido possedute da ciascuna sono uguali. Quindi il fluido ceduto dall'una $\Delta Q_1 = c_1 m_1 \Delta T_1$ dev'essere acquistato dall'altra la cui quantità di fluido aumenta di $\Delta Q_2 = c_2 m_2 \Delta T_2$. La quantità di fluido che passa dall'una all'altra sostanza ΔQ_i la possiamo chiamare **calore** che è una **grandezza fisica** perché misurabile: basta misurare la quantità di sostanza m_i , conoscere il tipo di sostanza per cui c_i è noto e misurare la variazione di temperatura ΔT_i .

Se il calore è una grandezza fisica gli si deve attribuire un'unità di misura: possiamo farlo scegliendo un processo da prendere come riferimento cui attribuire un valore convenzionale e arbitrario. Come riferimento possiamo prendere l'acqua: diremo che il calore si misura in **calorie** (cal) e che una caloria (1 cal) è la quantità di calore necessaria a innalzare la temperatura di 1 g d'acqua di 1 °C¹. In questo modo, dall'equazione

$$\Delta Q = cm\Delta T \quad (4.10)$$

ricaviamo che il coefficiente c non può essere adimensionale, come sembrava inizialmente, ma deve avere le dimensioni di

$$[c] = \frac{[\Delta Q]}{[m\Delta T]} \quad (4.11)$$

cioè di una quantità di calore per unità di massa e di variazione di temperatura. Di conseguenza c , che chiameremo **calore specifico** di una data sostanza, si misura in calorie per grado per grammo (cal/g·°C). Al prodotto $C = cm$ si dà il nome di

¹A dir la verità dovremmo essere piú precisi nel descrivere questo processo, ma il principio resta valido.

capacità termica in analogia alla capacità di un recipiente: tanto maggiore è questo prodotto, infatti, tanto maggiore è la quantità di calore che un corpo può immagazzinare.

4.2 Trasporto del calore

Quando il calore passa da un corpo piú caldo a uno piú freddo non lo fa istantaneamente. Prima che si raggiunga l'equilibrio termico occorre un certo tempo. Tra la fase iniziale in cui i due corpi hanno temperatura diversa e quella finale in cui i corpi sono tutti alla stessa temperatura ce n'è una di **non equilibrio** durante la quale avviene il passaggio di calore.

La quantità di calore che passa dal corpo caldo a quello freddo è evidentemente proporzionale al tempo durante il quale avviene il contatto: piú è lungo questo tempo piú calore passa. Se s'interrompe il passaggio prima che si raggiunga l'equilibrio una certa quantità di calore è comunque passata e le temperature dei corpi a contatto si modificano comunque, anche se non raggiungono una condizione di equilibrio. In prima approssimazione possiamo dire che il calore ΔQ che passa in un tempo Δt quando due corpi sono posti a contatto l'uno con l'altro è proporzionale a questo tempo:

$$\Delta Q \propto \Delta t. \quad (4.12)$$

Facendo qualche misura, anche grossolana, si capisce facilmente che la quantità di calore che transita da un corpo all'altro dipende dalla differenza di temperatura ΔT tra i due corpi: se si mette ghiaccio in acqua calda, questo si scioglie molto prima che se lo s'immerge in acqua appena uscita da un frigorifero. Possiamo quindi scrivere che

$$\frac{\Delta Q}{\Delta t} \propto \Delta T. \quad (4.13)$$

che significa che la **rapidità** con la quale il calore si trasferisce da un corpo all'altro dipende dalla differenza di temperatura tra i corpi. In questo caso ΔQ è la quantità di calore sottratta all'acqua e $\Delta T = T_{acqua} - T_{ghiaccio}$ la differenza di temperatura,

che è positiva. Se ΔQ è il calore **sottratto** dev'essere negativo perciò scriveremo più propriamente che

$$\frac{\Delta Q}{\Delta t} \propto -\Delta T. \quad (4.14)$$

Notate che se ΔQ rappresenta il calore acquistato dal ghiaccio (positivo), $\Delta T = T_{ghiaccio} - T_{acqua} < 0$ e il segno $-$ fa tornare i conti. È anche evidente che il calore che si può sottrarre a un corpo caldo dipende dal tipo di materiale di cui sono fatti i corpi. Se acquistate gelato in vaschetta, ve lo mettono in un recipiente di polistirolo e non in uno di metallo. Infatti il calore si trasmette con difficoltà dal polistirolo al gelato, mentre passa con relativa facilità dal metallo al gelato. Anche se ponete un cubetto di ghiaccio su una tavola di legno osserverete che si scioglie più lentamente di quanto faccia se poggiato su una lastra di metallo. Evidentemente è più facile per il calore passare dal metallo al ghiaccio piuttosto che dal legno al ghiaccio. Quest'osservazione potrebbe anche spiegare il fallimento del nostro **primo esperimento** che consisteva nel toccare la superficie di legno e i piedi in metallo di un banco: la sensazione che si provava era di freddo toccando il metallo e di caldo toccando il legno. Ora si comprende come questa sensazione non sia legata alla temperatura (che per quanto sopra dev'essere la stessa per legno e metallo), ma alla rapidità con la quale fluisce il calore dal nostro corpo alla superficie toccata, che è a temperatura più bassa.

Possiamo caratterizzare questo comportamento introducendo un coefficiente k che chiameremo **conducibilità termica** caratteristico di ogni materiale, e scriveremo che

$$\frac{\Delta Q}{\Delta t} \propto -k\Delta T. \quad (4.15)$$

Una volta che il calore ha abbandonato la superficie del corpo caldo, penetra in quello freddo attraverso la superficie di contatto ed è quindi ragionevole aspettarsi che la quantità di calore che passa per unità di tempo sia tanto più grande quando più è grande l'area A di questa superficie. Non ci vuo-

le molto a fare un esperimento di questo tipo per verificarlo, per cui dovrà anche essere che

$$\frac{\Delta Q}{\Delta t} \propto -kA\Delta T. \quad (4.16)$$

Del resto, è anche chiaro che nella fase di riscaldamento, la temperatura del corpo freddo comincia a cambiare vicino alla superficie di contatto col corpo caldo e solo col passare del tempo questo calore si diffonde a tutto il corpo. La rapidità con cui questo avviene dipende da quanto è esteso il corpo nella direzione perpendicolare alla superficie di contatto. Per tornare all'esempio del gelato, se vi consegnassero un chilo di gelato in una vaschetta di polistirolo dello spessore paragonabile a quello di un bicchierino per caffè usa e getta, probabilmente arriverebbe a casa sciolto. La vaschetta, per essere efficace, deve possedere uno spessore Δx ragionevolmente grande. Più è grande Δx e minore è la quantità di calore trasferita, perciò scriveremo che

$$\frac{\Delta Q}{\Delta t} \propto -kA \frac{\Delta T}{\Delta x}. \quad (4.17)$$

Il coefficiente di proporzionalità possiamo sceglierlo in maniera arbitraria se provvediamo a definire k con le opportune unità e naturalmente la scelta più conveniente è 1, quindi

$$\frac{\Delta Q}{\Delta t} = kA \frac{\Delta T}{\Delta x}. \quad (4.18)$$

Con questa scelta le dimensioni di k sono

$$[k] = \left[\frac{\Delta Q}{\Delta t} \frac{\Delta x}{\Delta T} \frac{1}{A} \right] = \left[\frac{QL}{tTL^2} \right] = \left[\frac{Q}{tTL} \right] \quad (4.19)$$

avendo indicato con $[t]$ le dimensioni fisiche del tempo e con T quelle della temperatura, per distinguerle. Le unità con le quali si misura la conducibilità termica sono dunque quelle di calorie al secondo per grado per metro ($\text{cal/s} \cdot ^\circ\text{C} \cdot \text{m}$). La modalità di trasferimento di calore che abbiamo analizzata è quella che si definisce **conduzione** che ha senso se si parla di passaggio di calore tra due solidi. Nel caso dei fluidi le caratteristiche che determinano il passaggio di calore da una sostanza all'altra sono diverse.

Studiamo un caso intermedio: quello di un solido in contatto termico con un fluido. L'esempio è quello del banco, la cui superficie e le cui gambe sono in contatto termico con l'aria circostante.

In questo caso, la prima parte delle osservazioni fatte nel caso della conduzione funziona ancora: la rapidità con la quale il calore passa da un corpo caldo a uno freddo dipende dalla differenza di temperatura e dalla superficie di contatto:

$$\frac{\Delta Q}{\Delta t} \propto -A\Delta T. \quad (4.20)$$

In questo caso per differenza di temperatura s'intende quella che esiste tra la superficie del solido e quella che il fluido possiede lontano da questa. Vicino alla superficie del solido, infatti, il fluido assumerà una temperatura simile. Il principio è usato nei **radiatori** (delle automobili, delle CPU dei computer, dei termosifoni) che sono costruiti in modo tale da avere la forma di molte *ali* in modo da massimizzare la superficie esposta all'aria. Anche in questo caso la quantità di calore che fluisce dipende da qualche caratteristica tipica del materiale che indichiamo con un coefficiente h

$$\frac{\Delta Q}{\Delta t} \propto -Ah\Delta T, \quad (4.21)$$

Ora non ha molto senso parlare di spessore di aria che il calore deve attraversare, almeno nei casi in cui l'aria non è contenuta essa stessa in un recipiente. In fondo il banco è in contatto con l'aria che si estende, al limite, a tutta quella contenuta nell'atmosfera terrestre. È come se fosse $\Delta x \simeq \infty$ e quindi non può essere lo spessore di aria a determinare la rapidità con la quale il calore si disperde. Dev'essere, anzi, qualche altra proprietà che prima non aveva senso considerare come la sua velocità: quando c'è vento abbiamo una percezione di freddo maggiore; l'aria calda proveniente da un asciugacapelli produce una maggiore sensazione di caldo se esce dall'elettrodomestico con velocità maggiore; lo sventolare di un ventaglio o la rotazione delle pale di un ventilatore producono sensazioni di fresco e sulle CPU dei computer si montano spesso ventole che smuovano l'aria circostante. Abbiamo già osservato che la sensazione

di caldo o di freddo non dipende dalla temperatura, ma dalla rapidità con la quale il calore fluisce via dal nostro corpo, perciò $\frac{\Delta Q}{\Delta t}$ deve dipendere dalla velocità v del fluido, ma non possiamo scrivere

$$\frac{\Delta Q}{\Delta t} \propto -Avh\Delta T, \quad (4.22)$$

perché il fluido che passa nelle vicinanze del solido non ha una velocità uniforme. Se dirigete il getto del vostro asciugacapelli parallelamente a un tavolo vedrete, disponendo sul tavolo coriandoli di carta, che lo strato di aria vicino al tavolo è molto veloce, ma la velocità diminuisce man mano che ci si allontana da questa.

L'unica cosa che possiamo dire, a questo livello, è che h deve dipendere dalle caratteristiche del moto del fluido (non solo dalla sua velocità, ma anche da quanto il moto è turbolento, da quanto il liquido è viscoso, etc.). Diremo allora che il passaggio di calore **per convezione**, come viene chiamato questo meccanismo, si può descrivere con l'equazione

$$\frac{\Delta Q}{\Delta t} = -Ah\Delta T, \quad (4.23)$$

avendo fatto le stesse considerazioni sulla scelta del coefficiente di proporzionalità che abbiamo fatto per la conduzione. Il valore di h lo si determina sperimentalmente ogni volta e al contrario di k non dipende solo dal tipo di materiale, ma anche dalle sue *condizioni* complessive.

Lo stesso meccanismo s'impiega per descrivere quel che avviene nello scambio di calore tra due fluidi. Per esempio, quando si prepara la colazione e il latte è troppo caldo, lo si agita col cucchiaino in modo che la superficie di contatto con l'aria s'incurvi, aumentando e favorendo la convezione del calore dal latte all'aria. Contemporaneamente, il moto turbolento del latte ne fa aumentare il valore di h il che favorisce la perdita di calore. Anche l'aria circostante si mette in moto, trascinata dalla superficie del latte che forma un vortice.

In definitiva la convezione consiste nel fatto che il trasporto del calore è affidato in gran parte al trasporto di materia: nel caso del contatto tra aria e solido, l'aria si sposta dai punti vicini al solido

a quelli lontano trasferendo lontano anche il calore perso dal solido; nel caso dei fluidi miscelati, come nel caso dei due liquidi, il moto del fluido in una certa misura *aiuta* il trasporto del calore. La parola **convezione** infatti è stata coniata dal latino **convehere** che vuol dire, appunto, **trasportare**.

Esiste un ulteriore meccanismo attraverso il quale il calore si propaga, che è chiamato **irraggiamento**. Accendendo un forno a legna e portandolo ad alta temperatura, avvicinandosi con le mani si percepisce il calore fluire fuori dal forno. Se però la bocca del forno è chiusa da un portello metallico, il calore che si percepisce è inferiore. La cosa potrebbe apparire *normale*, ma non è così: nel nostro modello quello che dovrebbe succedere è che inizialmente, appena chiuso il portello, il calore del forno è trasferito a quest'ultimo per convezione dal forno al metallo e quindi il metallo raggiunge relativamente presto la temperatura di equilibrio del forno. A questo punto dovrebbe cominciare a perdere calore esattamente come farebbe il forno se non ci fosse il portello. Ma non è così. Evidentemente una parte del calore riesce a sfuggire dal forno con la luce prodotta dalle fiamme. Coprendo quella luce si riduce l'effetto.

Tra l'altro il calore deve poter fluire anche in assenza di un mezzo che lo trasmetta, anche se con un'efficienza minore rispetto al caso della convezione o della conduzione. In effetti sulla Terra arriva il calore del Sole e non c'è praticamente alcun mezzo attraverso cui possa avvenire la conduzione o la convezione. Sole e Terra sono sostanzialmente separati dal vuoto. Misurando la quantità di calore che fluisce per unità di tempo per irraggiamento da corpi a diversa temperatura si trova che

$$\frac{\Delta Q}{\Delta t} = \sigma \varepsilon T^4 \quad (4.24)$$

dove T è la temperatura del corpo radiante. σ è una costante detta **costante di Stefan–Boltzmann** che vale $\sigma = 1.35 \times 10^{-8} \text{ cal}/(\text{s} \cdot ^\circ\text{C}^4 \text{m}^2)$. ε è un numero (adimensionale) $0 \leq \varepsilon \leq 1$ che dipende dalle caratteristiche della superficie del corpo radiante: 1 se è nera, 0 se bianca o a specchio.

4.3 Falsificare una teoria

Non ci vuole molto a capire che la teoria sviluppata due paragrafi sopra è **falsa**, nell'accezione che dà **Karl Popper** a questo termine (si veda, ad esempio, [5]). Con questo termine non s'intende una teoria sviluppata attraverso la *truffa* o avendo deliberatamente trascurato d'includere alcune osservazioni per corroborare il proprio punto di vista: significa che, sebbene la teoria spieghi le osservazioni dalle quali è stata dedotta, non riesce a spiegare osservazioni ulteriori che dovrebbero rientrare nell'ambito di questa. In sostanza le **predizioni** di questa teoria relativamente ai fenomeni che hanno portato alla sua formulazione sono **vere**, ma se s'includono altri fenomeni non siamo più in grado di spiegarli. Un altro filosofo della scienza, **Imre Lakatos** (vedi lo stesso libro di Giorello in [5]), direbbe che la teoria del calorico non possiede **euristica positiva**.

Vediamo quali sono i fenomeni che la teoria del calorico non riesce a spiegare. Basta fare due semplici esperimenti: il primo consiste nell'esporre un bicchier d'acqua al Sole. Come tutti sanno l'acqua si scalda. Il secondo esperimento consiste nel prendere un pezzo di carta vetrata e strofinare, con questa, un pezzo di metallo o di qualunque altro materiale²: in questo caso sia il metallo che la carta vetrata si scaldano³.

Se il primo esperimento si può spiegare ammettendo che la luce trasporti calore (ma della luce non sapremmo misurare la temperatura), il secondo è del tutto inspiegabile. Se davvero il calore fosse un fluido contenuto nei corpi, questo potrebbe passare da un corpo all'altro, ma come fa a passare in entrambi quando due corpi sono sfregati l'uno contro l'altro? Bisognerebbe ammettere che questo fluido si **genera** in qualche maniera quando i due corpi vengono a contatto, ma bisognerebbe capire come mai non basta il semplice contatto per produrre questo

²Si possono pensare molte varianti di questo esperimento, come quella che consiste nel fare un buco con un trapano, o anche fare esperimenti più semplici come quello di sfregarsi le mani.

³Nel caso in cui si fa un buco col trapano si scaldano sia la punta del trapano che l'oggetto forato.

fluido, ma occorre uno sfregamento. In effetti basta si verifichi una qualche forma di urto per generare calore perché, ad esempio, se si picchia forte con un martello la testa d'un chiodo questa si scalda! E lo sfregamento si può interpretare come una successione di piccolissimi urti tra le microscopiche sporgenze presenti sulla superficie dei corpi. In definitiva è necessario, per produrre calore, che due corpi siano in moto relativo tra loro e che qualcosa si opponga a questo moto.

Sebbene si possa pensare alla luce come formata da un flusso di corpuscoli che colpendo gli oggetti illuminati li riscaldino allo stesso modo di come fa un martello che colpisce un chiodo, altri esperimenti dimostrano che l'interpretazione di questo fenomeno non è così semplice. Se per esempio si collegano i due poli di una pila elettrica con un filo metallico, quest'ultimo si scalda, senza che apparentemente ci sia alcunché che si muova. Anche se si versano certi sali in acqua, quando questi si sciolgono provocano l'innalzamento o l'abbassamento della temperatura del liquido (questo sistema si usa per produrre le bustine di ghiaccio istantaneo usate per lenire i dolori da contusioni o le bibite calde espresse, nelle quali basta pigiare sul fondo del recipiente per rompere un contenitore stagno che libera dei sali che fanno salire la temperatura del liquido contenuto sul fondo del recipiente). Ulteriori fenomeni legati ai passaggi di calore, apparentemente inspiegabili alla luce della teoria del calorico, sono i cambiamenti di stato: sotto una certa temperatura l'acqua diventa ghiaccio (solidifica) e sopra i 100 °C bolle diventando vapore.

Durante il processo di solidificazione (o fusione) e quello di ebollizione la temperatura dell'acqua non cambia come si può verificare con un termometro, la cui lunghezza non cambia fino a quando sono presenti contemporaneamente la fase liquida e quella solida nel caso della solidificazione (fusione) o quella liquida e quella gassosa. Solo quando l'acqua è diventata tutta ghiaccio, la temperatura di quest'ultimo continua a scendere e solo quando tutta l'acqua s'è trasformata in vapore la sua temperatura ricomincia a salire.

Che in questi casi ci sia un passaggio di calore è dimostrato dal fatto che la temperatura di equili-

brio di un sistema in cui si mescolino le fasi cambia al cambiare della quantità di sostanza che cambia stato. Per esempio, mescolando m_g g di ghiaccio con m_a g di acqua, la temperatura T dell'acqua che si raggiunge se il ghiaccio si scioglie completamente è proporzionale alla quantità di ghiaccio sciolta e inversamente proporzionale alla massa d'acqua. In formule possiamo scrivere

$$T = T_i - \frac{\lambda m_g}{cm_a} \quad (4.25)$$

dove T_i è la temperatura iniziale dell'acqua. Questo significa che per sciogliere m_g g di ghiaccio serve una quantità di calore ΔQ pari a

$$\Delta Q = \lambda m_g \quad (4.26)$$

e la costante λ prende il nome di **calore latente di fusione** del ghiaccio e vale $\lambda \simeq 80 \text{ cal/g}$ ⁴. In maniera del tutto analoga, per far evaporare una quantità d'acqua m_v occorre una quantità di calore

$$\Delta Q = L m_v \quad (4.27)$$

dove L , detto **calore latente di vaporizzazione** dell'acqua, vale $L \simeq 540 \text{ cal/g}$.

Gli esperimenti descritti, quanto meno, ci dicono che se il fluido calorico esiste, non può essere qualcosa di **contenuto** nei corpi, ma si deve poter **generare**. Cosa lo generi e di cosa effettivamente sia fatto non è immediato da spiegare. Di sicuro però abbiamo diverse osservazioni importanti di cui tenere conto:

- la luce trasporta calore;
- il moto deve aver a che fare col calore;
- il calore si può generare nelle reazioni chimiche;
- il calore può provocare il cambiamento di stato delle sostanze;
- il passaggio di corrente elettrica scalda i fili metallici.

⁴L'acqua ha perso una quantità di calore $\Delta Q = cm_a \Delta T = \lambda m_g$ che non può che essere servita per sciogliere il ghiaccio.

Se vogliamo capire qualcosa di piú è necessario studiare prima di tutto la **luce** e il **moto**. Quindi dovremo cercare di capire qualcosa di piú di come sono fatti i corpi delle diverse sostanze. Infine bisognerà capire cosa c'entra l'elettricità col calore.

Ottica geometrica

La luce è qualcosa che si **propaga** nei mezzi trasparenti come l'aria, il vetro, certe materie plastiche, etc.. Il fatto che la luce non sia qualcosa di *presente* in un determinato luogo, ma che *passa* da quel luogo, è confermato dal fatto che, abbassando le serrande delle finestre, all'interno di una stanza si fa il buio. Questo vuol dire che la luce che prima era presente nella stanza proveniva dall'esterno e che, a causa dell'ostacolo rappresentato dalle serrande, non può più entrare. La luce che era entrata dalla finestra viene in qualche modo **assorbita** dai corpi. Se così non fosse, poiché una certa quantità di luce era entrata nella stanza prima di chiudere le serrande, dovrebbe continuare a rimanere in essa anche con la chiusura delle serrande; al massimo ci possiamo aspettare una lieve diminuzione dovuta alla possibilità che una parte di essa esca dalle finestre prima che siano completamente chiuse. È presumibile che questo assorbimento dia luogo al **risaldamento** che osserviamo quando esponiamo un corpo alla luce.

Prima di cercare di capire cosa sia la luce è necessario studiarne il comportamento, perché possiamo dedurre la natura delle cose soltanto dai fatti sperimentali. La prima cosa da studiare è, evidentemente, il modo in cui la luce si propaga. Lo studio della propagazione della luce si chiama **ottica geometrica**.

Sempre osservando la luce che passa attraverso una finestra possiamo renderci conto che la luce si propaga in linea retta. Anche se ci sono dei casi in cui questo sembra non avvenire, come nelle **fibre ottiche** nelle quali la luce sembra propagarsi seguendo le curve prodotte dalle fibre, ma in questo caso la luce è in qualche maniera *costretta* a restare

all'interno della fibra per uscirne solo alle estremità.

Possiamo schematizzare la propagazione della luce immaginando che sia composta di **raggi** rappresentabili come segmenti che hanno una **direzione** (che definisce la retta su cui giacciono) e un **verso** (che rappresenta in quale dei due possibili sensi la retta è percorsa dal raggio).

Una buona approssimazione di raggio si può ottenere preparando uno schermo opaco (che quindi non lascia passare la luce) con un foro al centro, che lascia passare solo uno stretto *pennello* di luce. Non ci vuole molto a rendersi conto che quest'approssimazione è valida solo per distanze relativamente piccole: il fascio in uscita dal foro infatti s'allarga sempre di più ed è più simile a un cono che a un fascio cilindrico di raggi che corrono nella stessa direzione.

La luce di un **laser** però si presenta abbastanza simile a quello che possiamo pensare sia un **raggio**: la luce che esce dal laser si propaga parallelamente a sé stessa e anche dopo aver percorso una distanza consistente il fascio di luce appare comunque stretto. Il laser, insomma, è un utile strumento per fare esperimenti di ottica geometrica.

5.1 Riflessione della luce

Inviando luce su una superficie lucida, come quella di uno specchio, quello che si osserva è che almeno una frazione di essa cambia direzione, in modo tale da formare un raggio che si muove sul piano formato dal raggio incidente e dalla perpendicolare allo specchio nel punto in cui incide e in modo da formare lo stesso angolo con la perpendicolare alla superficie dello specchio. In formule

$$\theta_{inc} = \theta_{rifl}. \quad (5.1)$$

Quest'uguaglianza esprime la **legge della riflessione**: l'angolo formato dal raggio incidente e la normale alla superficie si chiama **angolo d'incidenza**, indicato con θ_{inc} , quello formato dal raggio riflesso e la normale alla superficie si chiama **angolo di riflessione**, indicato con θ_{rifl} . La legge si esprime dicendo che l'angolo di riflessione è uguale all'angolo d'incidenza.

La percentuale di luce riflessa, rispetto a quella incidente, dipende dalle caratteristiche della superficie e dall'angolo d'incidenza. Le caratteristiche della superficie sono condensate in un parametro detto **riflettività** o **riflettanza** ρ dello specchio, con $0 \leq \rho \leq 1$. Per uno specchio perfetto $\rho = 1$.

Attraverso lo specchio si vede un'immagine di ciò che sta davanti. La cosa si spiega abbastanza facilmente se consideriamo il meccanismo che ci permette di osservare gli oggetti. Innanzi tutto è chiaro che per vedere qualcosa è indispensabile che sia illuminata. Quello che succede, evidentemente, è che la luce che finisce sull'oggetto è **diffusa**¹ in tutte le direzioni. La luce diffusa è chiaramente percepita dai nostri occhi che ricostruiscono nel cervello l'immagine degli oggetti da cui la luce proviene. Un modo per capire cosa succede quando la luce diffusa proveniente da un punto incontra uno specchio consiste nel considerare due o più raggi provenienti da un punto qualunque dell'oggetto e seguirne la traiettoria.

Schematizziamo un oggetto posto dinanzi a uno specchio come una freccia, come in Fig. 5.1. Dal punto estremo dell'oggetto, indicato con A , provengono infiniti raggi di luce che si propagano in tutte le direzioni. Uno di essi è quello rappresentato in rosso che incide perpendicolarmente allo specchio (in blu) posto a una certa distanza dall'oggetto. Questo raggio incide sullo specchio nel punto B e per la legge della riflessione torna indietro nella stessa direzione,

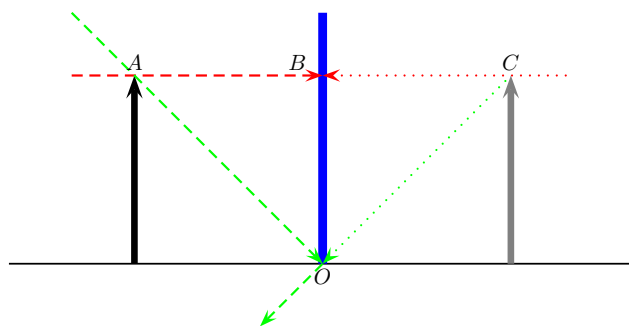


Figura 5.1 Costruzione dell'immagine prodotta da uno specchio piano.

solo con verso opposto. Il raggio verde, che si propaga fino al punto O , invece è riflesso verso il basso in modo da formare un angolo di riflessione uguale a quello d'incidenza. Dallo specchio, i raggi riflessi appaiono provenire entrambi da uno stesso punto indicato con la lettera C nella figura. Di conseguenza, così come il nostro occhio interpreta l'insieme dei raggi provenienti da A come l'estremo della freccia, interpreterà i raggi provenienti da C esattamente nello stesso modo e quindi vedrà l'immagine della freccia **dentro** lo specchio, a una distanza da questo uguale a quella che c'è tra oggetto e specchio. L'immagine prodotta dallo specchio è rappresentata in grigio nella figura.

Se lo specchio non è piano l'immagine appare deformata. Evidentemente lo specchio può avere qualunque forma, ma è chiaro che possiamo studiare solo forme relativamente semplici. Quindi cominciamo con una forma regolare come quella sferica. Uno **specchio sferico** è una superficie riflettente che ha la forma di una calotta sferica. Naturalmente la calotta può essere riflettente sia sulla superficie esterna (in questo caso lo specchio si dice **convesso**) sia su quella interna (e in tal caso si dice **concavo**).

L'effetto prodotto dai due tipi di specchi è diverso: gli specchi cosmetici, per esempio, sono specchi concavi, mentre quelli stradali o da sorveglianza sono convessi. I primi hanno la proprietà di ingrandire

¹Non riflessa, perché solo la luce che si propaga con un angolo pari a quello d'incidenza si può dire riflessa: quella diffusa è luce che si propaga in direzioni non necessariamente correlate con la direzione d'incidenza.



Figura 5.2 La pallina di un albero di Natale con la superficie riflettente è uno **specchio sferico convesso**. Sulla sue superficie le immagini appaiono più piccole di quanto siano in realtà (foto di Christopher Bachner).

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 5.3 Un raggio di luce (in rosso) che incide parallelamente all'asse ottico di uno specchio sferico è riflesso in modo tale da incrociare l'asse ottico nel punto F detto **fuoco** dello specchio [<http://youtu.be/gnvRSvCIS3I>].

molto gli oggetti riflessi (si usano, infatti, per truccarsi o per radersi), ma solo se si sta abbastanza vicini. Se osservate uno specchio cosmetico (concavo) da lontano vedrete che l'immagine riflessa è rovesciata rispetto agli oggetti presenti nella stanza. Al contrario, gli specchi stradali o da sorveglianza (convessi) producono sempre immagini dritte e rimpicciolite degli oggetti. Cerchiamo di capire perché.

Cominciamo con l'analizzare un caso semplice, consistente nell'inviare sulla superficie di uno specchio, un raggio parallelo alla direzione di quello che chiamiamo **asse ottico**. L'asse ottico è la retta passante per il vertice della calotta sferica e perpendicolare a questa. Quando il raggio incontra la superficie

dello specchio (Filmato 5.3) è riflesso in maniera tale che l'angolo $\widehat{ORF}$ (in verde) formato dal raggio riflesso e la normale nel punto d'incidenza sia uguale a quello definito da quest'ultima e il raggio incidente (in rosso col vertice in R). Poiché i raggi di una sfera sono sempre perpendicolari alla sua superficie, la retta normale al punto d'incidenza è diretta come il raggio della sfera $\overline{OR}$, lungo r . L'asse ottico (tratteggiato) è parallelo al raggio incidente, quindi l'angolo $\widehat{ROF}$ è uguale a quello d'incidenza e, di conseguenza, è anche uguale a quello di riflessione. Il triangolo OFR allora è isoscele e i lati $\overline{OF}$ e $\overline{FR}$ sono uguali. Se si riduce la distanza tra il raggio incidente e l'asse ottico le relazioni sopra descritte continuano a valere: in particolare la distanza tra il centro della superficie sferica e il punto in cui il raggio riflesso incontra l'asse ottico è sempre uguale alla distanza tra questo punto e il punto in cui incide la luce. Man mano che il raggio incidente si avvicina all'asse ottico il triangolo si *schiaccia* sempre più e gli angoli alla base diventano sempre più piccoli fino a quando, nel momento in cui la distanza dall'asse ottico vale zero, il triangolo degenera in un segmento la cui base è lunga $\overline{OR} = r$, gli angoli uguali sono nulli e i due lati uguali valgono ciascuno $r/2$. Il Filmato 5.3 mostra cosa succede al variare della distanza del raggio incidente dall'asse ottico: il raggio riflesso incontra l'asse ottico in un punto, che chiameremo **fuoco**, in blu, che cambia al cambiare della distanza del raggio incidente dall'asse ottico. Per angoli anche relativamente grandi, la variazione della posizione del fuoco, tuttavia, è piuttosto piccola e, in particolare per distanze molto piccole, vale la relazione

$$f \simeq \frac{r}{2} \quad (5.2)$$

avendo chiamato f la distanza $\overline{OF}$. Se quindi ci limitiamo a specchi le cui dimensioni siano abbastanza piccole in confronto ai loro raggi di curvatura, possiamo stabilire che il fuoco di uno specchio sferico si trova a una distanza dal vertice pari a metà del raggio di curvatura e che tutti i raggi di luce che incidono parallelamente all'asse ottico devono passare per il fuoco. Un raggio di luce che invece incide

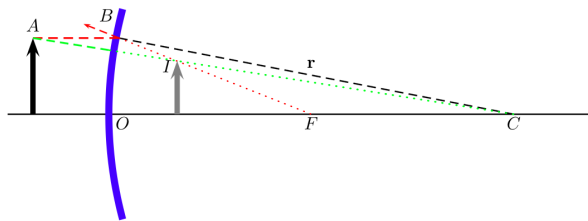


Figura 5.4 Costruzione dell'immagine di un oggetto prodotta da uno specchio sferico convesso. Una costruzione **animata** dell'immagine si può vedere all'indirizzo <http://tube.geogebra.org/student/m682967>.

sullo specchio viaggiando lungo uno dei suoi raggi evidentemente è riflesso all'indietro.

Con queste due regole possiamo facilmente costruire un'immagine di qualsiasi oggetto: cominciamo da uno specchio convesso, che ha un comportamento univoco. Facciamo riferimento alla Figura 5.4: rappresentiamo l'oggetto come una freccia (la punta della quale è marcata con la lettera A nella figura). Un raggio di luce proveniente dalla punta della freccia², che si dirige verso lo specchio parallelamente all'asse ottico, incide su questo in B . Qui è riflesso in modo tale che l'angolo di riflessione sia uguale a quello d'incidenza ed è facile convincersi che il prolungamento del raggio riflesso (in rosso tratteggiato) si dirige verso il fuoco dello specchio. La luce è riflessa dall'altro lato, quindi chi osserva la superficie dello specchio vede la luce arrivare da un punto che si trova lungo il prolungamento del raggio, che ha la proprietà di passare per il fuoco. D'altra parte un raggio che provenga dallo stesso punto dell'oggetto che incide perpendicolarmente alla superficie (in verde nella figura) è riflesso all'indietro e appare provenire dal centro dello specchio³. Il raggio

²Quando fate questi disegni usate sempre un righello: la precisione è essenziale per determinare correttamente quel che accade; basta una linea leggermente storta per cambiare completamente il risultato!

³In alternativa si può usare un raggio (o un suo prolungamento) che passi attraverso il fuoco dello specchio. Questo

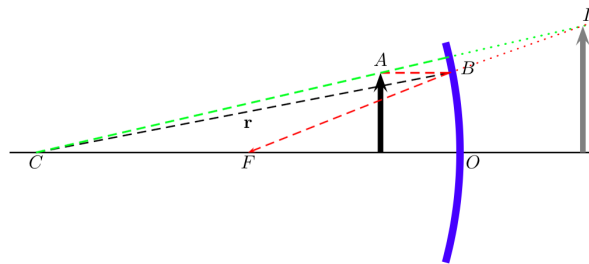


Figura 5.5 Costruzione dell'immagine di un oggetto prodotta da uno specchio sferico concavo quando l'oggetto si trova a una distanza dallo specchio $d < f$.

verde e il raggio rosso hanno il punto I in comune: quindi chi osserva lo specchio vede tutti i raggi provenire da questo punto, come nel caso dell'oggetto reale. Di conseguenza la punta dell'oggetto sembra trovarsi nel punto I e dunque l'oggetto appare come la freccia in grigio a destra dello specchio. Chi guarda lo specchio dunque vede gli oggetti che sono al di qua come se si trovassero dentro lo specchio e rimpiccioliti.

Se invece lo specchio è concavo vediamo tre fenomeni diversi, secondo la distanza alla quale si trova l'oggetto dallo specchio. In effetti, mentre per lo specchio convesso tutte le distanze sono equivalenti, nel caso degli specchi concavi l'oggetto si può trovare a una distanza d dallo specchio che è inferiore alla distanza focale f , superiore a questa ma inferiore a r e maggiore di r . Esistono dunque tre condizioni che in principio possono essere diverse e che vanno per questo analizzate separatamente.

Cominciamo dal caso $d < f$, facendo riferimento alla Figura 5.5. Tracciamo sempre un raggio proveniente dalla punta dell'oggetto che viaggia verso lo specchio parallelamente all'asse ottico (in rosso). Nel punto B è riflesso in modo da dirigersi nel fuoco F . Un raggio che invece si muove nella direzione perpendicolare allo specchio (in verde) è riflesso all'indietro lungo la stessa direzione. Chi si trova al di qua dello specchio dunque vede provenire tutti i

raggio è riflesso dallo specchio parallelamente all'asse ottico. Il suo prolungamento passa per il punto I .

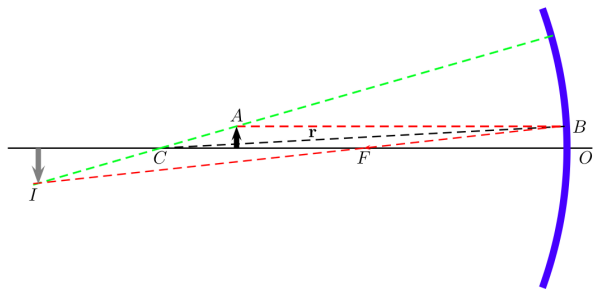


Figura 5.6 Costruzione dell'immagine di un oggetto prodotta da uno specchio sferico concavo quando l'oggetto si trova a una distanza dallo specchio $f < d < r$.

raggi che partono dallo specchio dal punto nel quale i prolungamenti dei raggi riflessi (linee puntinate) s'incrociano. Questo punto, marcato con I si trova al di là dello specchio e l'immagine appare ingrandita rispetto all'originale.

Consideriamo ora il caso $f < d < r$, riportato nella Figura 5.6. Il raggio che si muove parallelamente all'asse ottico, proveniente dalla punta dell'oggetto, è riflesso nel fuoco e prosegue verso il punto I , dove s'incrocia con quello (verde) riflesso all'indietro perché incide perpendicolarmente alla superficie dello specchio. Si vede chiaramente che l'immagine dell'oggetto appare ingrandita e rovesciata. Osserviamo che adesso la situazione è un po' diversa da prima! Nei primi casi analizzati l'immagine si trovava *all'interno* dello specchio e i raggi di luce sembravano provenire da essa, perciò i nostri occhi potevano percepire l'immagine come se si trovasse proprio dentro lo specchio. In questo caso l'immagine si forma dallo stesso lato nel quale si trova l'oggetto. I raggi di luce non partono dall'immagine, ma convergono in essa! Se mettiamo uno schermo nel punto in cui convergono i raggi, su questo schermo si forma un'immagine dell'oggetto davanti allo specchio. Per questa ragione l'immagine si dice **reale**, mentre nel caso dell'immagine al di là dello specchio si parla di immagine **virtuale**. Le immagini virtuali sono tali perché i raggi di luce *appaiono* provenire da essa, ma di fatto la luce proviene da un altro punto. Le

immagini reali invece si possono osservare su uno schermo posto nel punto in cui convergono i raggi. Quando guardiamo la superficie di uno specchio concavo, nel caso delle immagini virtuali abbiamo l'impressione che l'immagine si trovi al di là dello specchio perché la luce sembra provenire da lì. Se gli oggetti si trovano a una distanza $d > f$ dallo specchio le immagini che si formano sono reali. Questo vuol dire che, guardando lo specchio, vediamo i raggi di luce provenire *realmente* da queste. L'immagine della freccia della Figura 5.6 si può osservare solo se ci troviamo a una distanza dallo specchio maggiore di quella alla quale si forma l'immagine (a sinistra della freccia grigia). Guardando in direzione dello specchio vedremo raggi di luce provenire dal punto I che riproducono fedelmente quelli che provengono dal punto A quindi vedremo un'immagine capovolta e ingrandita dell'oggetto e non dietro lo specchio, ma davanti. Il caso $d > r$ è lasciato come esercizio.

Esercizio 5.1 *Uno specchio sferico*

Costruisci l'immagine di un oggetto prodotta da uno specchio sferico concavo nel caso in cui l'oggetto si trova a una distanza dallo specchio maggiore del suo raggio di curvatura. Disegna l'oggetto come una freccia verticale, quindi traccia una retta parallela all'asse ottico che parte dalla punta della freccia e incontra lo specchio. Usa un righello per fare il disegno: una minima inclinazione della retta rispetto all'asse ottico può produrre risultati sbagliati. Dal punto in cui il raggio incontra lo specchio traccia una retta che passa per il fuoco dello specchio, che si trova a una distanza pari a $f \simeq r/2$ dal suo vertice. Quindi disegna una retta che passa per la punta della freccia e per il centro dello specchio. La luce che viaggia su questa retta è riflessa all'indietro. Dove s'incontrano le rette? L'immagine è reale o virtuale? Dritta o capovolta? Ingrandita o rimpicciolita? Si può vedere l'immagine proiettata su uno schermo? Confronta i risultati con quanto puoi vedere all'indirizzo <http://tube.geogebra.org/student/682961>.

5.2 Una prima interpretazione

Quel che si evince analizzando il funzionamento degli specchi è che la luce, quando colpisce uno specchio, si comporta come farebbe un pallone che *rimbalza* colpendo un muro o la biglia di un biliardo che urta la sponda. Affinché la biglia rimbalzi è necessario che la sponda abbia certe caratteristiche di rigidità e di elasticità: una biglia non rimbalza su un birillo o su un'altra biglia. Allo stesso modo non tutte le superfici sono in grado di produrre una **riflessione**.

Sembrerebbe dunque ragionevole considerare la luce composta di un flusso di **corpuscoli** che devono essere molto piccoli visto che non riusciamo a **risolverli**, cioè a vederli distintamente. I corpuscoli devono di norma propagarsi su linee rette (sicuramente a velocità molto alte, al limite infinite). I fenomeni che abbiamo osservato finora sono spiegabili



Figura 5.7 Un **mirascopio** è un dispositivo che produce immagini che sembrano reali, ma non lo sono affatto. La piccola rana che si vede in questa figura si trova all'interno dello strumento, ma si ha l'illusione che sia al di sopra di esso.

in termini di corpuscoli che rimbalzano quando urtano certi materiali o che sono assorbiti colpendone altri. Visto che al buio non vediamo nulla, evidentemente i nostri occhi sono sensibili alla luce che proviene dagli oggetti che vediamo i quali non devono poterla emettere autonomamente, perché altrimenti non sarebbe necessario accendere una lampada per vederli al chiuso. Quello che potrebbe succedere è che dalla lampada sono emessi questi corpuscoli che urtano la superficie degli oggetti. Una parte di questi corpuscoli è assorbita dagli oggetti (questo, tra l'altro, potrebbe provocarne il **riscaldamento**), un'altra parte dev'essere **diffusa**, cioè deviata in tutte le possibili direzioni. In effetti riusciamo a vedere gli oggetti illuminati da qualunque direzione si guardi, mentre le immagini riflesse si vedono solo se l'angolo formato tra lo specchio, il raggio incidente e quello riflesso assume un particolare valore. È possibile che la superficie degli oggetti non sia liscia, anche qualora appaia come tale, ma più o meno scabra, perciò quando i corpuscoli luminosi la colpiscono possono rimbalzare in direzioni diverse, secondo l'angolo con il quale urtano la superficie.

Il Mirascopio

Un **mirascopio** è uno strumento ottico che produce un effetto stupefacente. È costituito di due specchi di forma parabolica posti l'uno sopra l'altro. Gli specchi parabolici si comportano come quelli sferici, ma al contrario di questi ultimi, hanno la proprietà di avere il fuoco in un punto preciso (ricordate che per gli specchi sferici i raggi paralleli all'asse ottico non convergono tutti in un solo punto, ma in una regione che, solamente se sufficientemente piccola, si può assimilare a un punto). Sullo specchio superiore è praticato un foro. Ponendo un oggetto all'interno del mirascopio, per esempio appoggiandolo sulla superficie dello specchio in basso, i raggi di luce provenienti da uno dei suoi punti sono riflessi da uno degli specchi, che li invia sull'altro. Quest'ultimo produce un'ulteriore riflessione che porta il raggio a uscire dalla parte superiore dello strumento dove si trova un foro. La geometria degli specchi è tale da far convergere tutti i raggi provenienti da un punto qualunque dell'oggetto posto all'interno in uno stesso punto al di fuori del microscopio. Il risultato è che i raggi di luce provenienti da un punto dell'oggetto appaiono provenire dall'esterno del microscopio creando l'illusione che l'oggetto sia realmente all'esterno.

All'indirizzo <https://www.geogebraTube.org/student/m682973> si trova una simulazione di un mirascopio fatta con **GEOGEBRA**. Piccoli mirascopi sono in vendita nei negozi di giocattoli o su Internet.

Una parte di questi entra nei nostri occhi e produce la sensazione visiva. Gli specchi devono essere superfici molto lisce (e in effetti lo sono): talmente lisce da far rimbalzare i raggi luminosi sempre nella stessa direzione.

I corpuscoli che formano la luce non devono essere tutti uguali tra loro. Devono esistere corpuscoli diversi che trasmettono sensazioni di colore diverse. È chiaramente la luce che l'illumina a fornire il colore agli oggetti: se ci chiudiamo in una stanza il-

luminata da una lampada rossa tutto il contenuto della stanza appare rosso. Il colore dunque non è una proprietà degli oggetti: semmai è una proprietà della luce diffusa dagli oggetti. Ma perché un oggetto illuminato con luce bianca appare colorato? Molto probabilmente i corpuscoli che formano la luce bianca sono una miscela di corpuscoli di vari colori: se s'illumina la copertina blu d'un libro, evidentemente i corpuscoli di colori diversi sono assorbiti dal libro, mentre quelli blu sono diffusi e li vediamo.

Per verificare quest'ipotesi è necessario fare un esperimento che la confermi. Un primo esperimento consiste nel far passare un raggio di luce bianca (per esempio quello prodotto dalla luce di un videoproiettore fatto passare attraverso una fenditura) attraverso un prisma di vetro. Curiosamente dalla parte opposta si forma un'immagine che ha la forma di un rettangolo colorato.

I colori sono gli stessi (e nella stessa sequenza) che si osservano nel caso del verificarsi di un **arcobaleno**. È molto probabile che il meccanismo che produce questi due fenomeni sia quindi lo stesso. È come se il prisma avesse la proprietà di *separare* le componenti del fascio bianco in componenti colorate. Il fenomeno è chiamato **dispersione** dai fisici. Se dalla luce bianca si ottiene luce colorata è chiaro che la luce bianca dev'essere il risultato del *sommarsi* degli effetti prodotti sul nostro occhio dai corpuscoli dei vari colori. A ulteriore conferma si può realizzare quello che si chiama un **disco di Newton**: disegnate su un disco tanti spicchi e colorate ciascuno di essi con tutti i colori dell'arcobaleno; poi fatelo ruotare rapidamente attorno al proprio asse (montandolo su un motorino elettrico, per esempio, o su un trapano). Il disco appare bianco! Sappiamo bene che sulla retina dei nostri occhi le immagini permangono per qualche frazione di secondo prima di svanire e che l'effetto si usa nel cinema in cui proiettando su uno schermo immagini leggermente diverse l'una dall'altra si ha l'impressione di vedere un movimento fluido. Il fatto che il disco appare bianco si può spiegare nello stesso modo: i corpuscoli colorati provenienti dai vari settori del disco giungono rapidamente uno dopo l'altro ai nostri occhi. Se

il tempo necessario a far giungere tutti i corpuscoli è inferiore a quello di permanenza della sensazione prodotta nel nostro cervello è come se all'occhio giungessero contemporaneamente corpuscoli di tutti i colori.

La nostra teoria sta funzionando egregiamente! Ma non bisogna mai accontentarsi dei primi successi.

5.3 La rifrazione

Quando la luce attraversa la superficie di separazione tra due mezzi trasparenti, come ad esempio l'aria e l'acqua, il suo cammino subisce una deviazione e si dice che il raggio è **rifratto**⁴. Basta immergere una cannuccia in un bicchier d'acqua per rendersi conto che la luce proveniente dalla parte immersa deve giungere ai nostri occhi con un angolo diverso da quella proveniente dalla parte emersa, perché la cannuccia appare spezzata o piegata. Per capire meglio il fenomeno conviene come al solito eseguire esperimenti in condizioni controllate. Per esempio, si può illuminare il fondo di un acquario con un raggio laser e marcare il punto raggiunto dalla luce con un pennarello. Se, senza spostare nulla, si versa acqua nell'acquario, il punto luminoso sul fondo si sposta e osservando il raggio diffuso dall'acqua si nota che, quando ne attraversa la superficie, cambia direzione.

Non resta che fare qualche misura per capire come cambia la direzione del raggio luminoso. Nel caso della riflessione avevamo scoperto che la direzione della luce riflessa era caratterizzata dall'angolo θ_{inc} formato tra il raggio incidente e la normale alla superficie dello specchio. Di certo anche in questo caso è così: cambiando l'inclinazione della luce incidente cambia l'angolo di rifrazione θ_{rifr} che quindi dev'essere funzione di θ_{inc} : $\theta_{rifr} = f(\theta_{inc})$.

Eeguire le misure non è difficile: è sufficiente disporre di un supporto per il laser che permetta di orientarlo con diverse inclinazioni rispetto alla verticale e di un acquario riempito d'acqua fino a un'al-

tezza h . Mettendo il laser in posizione verticale si osserva che la luce del laser non è deviata quando penetra nell'acqua. Man mano che l'angolo d'inclinazione aumenta, lo *spot* luminoso si allontana dalla verticale. Se sul fondo della vasca abbiamo incollato un righello possiamo misurare di quanto si sposta in funzione dell'angolo d'incidenza e così ricavare l'angolo di rifrazione come

$$\theta_{rifr} = \arctan \frac{d}{h} \quad (5.3)$$

dove d è la distanza raggiunta dal punto luminoso sul fondo rispetto al punto illuminato quando il laser è in posizione verticale. Misuriamo dunque il valore di θ_{rifr} in funzione di θ_{inc} . Osserviamo che se l'angolo è piccolo $\arctan \theta_{rifr} \simeq \theta_{rifr}$, quindi $\theta_{rifr} \simeq d/h$ e l'errore che si commette nel determinare l'angolo è l'errore sulla misura del rapporto. Per valutarlo facciamo così: l'errore che si commette nel misurare d è σ_d e quindi d è nota con una precisione di

$$\varepsilon_d = \frac{\sigma_d}{d} \quad (5.4)$$

che si chiama **errore relativo** della misura. Noto l'errore relativo di una grandezza fisica la sua entità si ricava semplicemente invertendo la relazione che definisce ε :

$$\sigma_d = d\varepsilon_d. \quad (5.5)$$

Analogamente la misura di h avrà una precisione di σ_h/h . L'errore relativo totale sarà la **somma (in quadratura)** degli errori sulle singole misure, quindi è

$$\varepsilon_\theta = \sqrt{\left(\frac{\sigma_d}{d}\right)^2 + \left(\frac{\sigma_h}{h}\right)^2}. \quad (5.6)$$

Se $h \simeq 10$ cm e $\sigma_h = \sigma_d \simeq 1$ mm, per un angolo dell'ordine dei 5° , $d \simeq 7$ mm e abbiamo quindi che

$$\varepsilon_\theta = \sqrt{\left(\frac{0.1}{0.7}\right)^2 + \left(\frac{0.1}{10}\right)^2} \simeq 0.14. \quad (5.7)$$

Conosceremo perciò l'angolo con una precisione del 14 %. Dato che il rapporto $d/h = 0.07$, la misura di questo rapporto si scrive

⁴Dal latino **refringere**: rompere, spezzare (è la stessa radice di **frangere**).

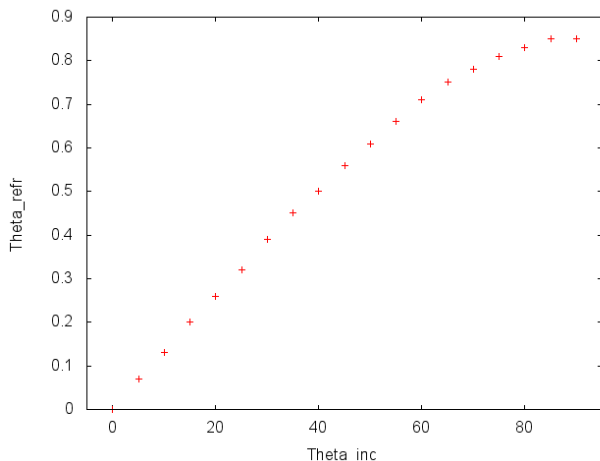


Figura 5.8 Grafico dell'angolo di rifrazione in acqua in funzione dell'angolo d'incidenza.

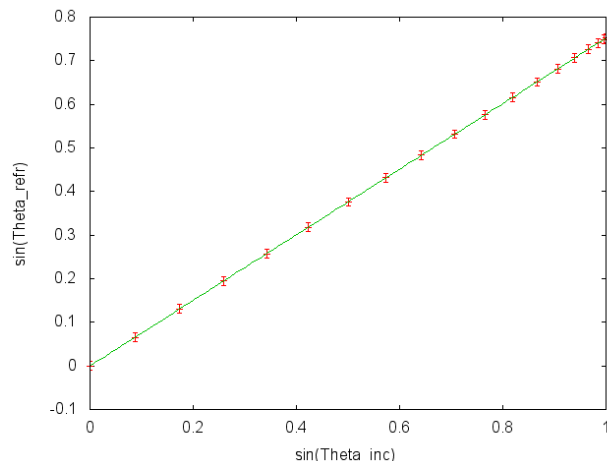


Figura 5.9 Grafico del seno dell'angolo di rifrazione in acqua in funzione del seno dell'angolo d'incidenza.

$$\theta_{rifr} \simeq \frac{d}{h} = 0.07 \pm 0.01, \quad (5.8)$$

visto che $0.14 \times \frac{0.7}{10} \simeq 0.01$. Naturalmente questa misura è in radianti. Se volessimo esprimerla in gradi bisognerebbe convertirla applicando l'opportuno fattore di conversione:

$$\theta_{rifr} [^\circ] = \theta_{rifr} [\text{rad}] \frac{180}{\pi} \quad (5.9)$$

da cui si ricava che

$$\theta_{rifr} \simeq (4.0 \pm 0.6)^\circ. \quad (5.10)$$

Nella Figura 5.8 riportiamo i valori degli angoli misurati con un dispositivo come quello descritto.

Dal grafico è chiarissimo che c'è una dipendenza di θ_{rifr} da θ_{inc} , ma che questa dipendenza non è lineare, se non approssimativamente nella prima parte del grafico. Questo significa che solo per angoli piccoli $\theta_{rifr} \simeq \alpha \theta_{inc}$. Se la dipendenza non è lineare allora la funzione che lega θ_{rifr} a θ_{inc} non può che essere una funzione trigonometrica. Ma in questo caso è improbabile che la relazione sia del tipo $\theta_{rifr} = f(\theta_{inc})$, perché almeno in certi casi le funzioni trigonometriche sono limitate, mentre θ_{rifr} no. È molto più probabile che la relazione che lega i due angoli sia del tipo

$$f(\theta_{rifr}) = g(\theta_{inc}). \quad (5.11)$$

L'ipotesi più semplice che possiamo fare è che la relazione sia del tipo

$$\sin \theta_{rifr} = \alpha \sin \theta_{inc}. \quad (5.12)$$

In effetti, con quest'assunzione, per angoli piccoli si ottiene che

$$\theta_{rifr} \simeq \alpha \theta_{inc}. \quad (5.13)$$

La Figura 5.9 mostra il seno di θ_{rifr} in funzione del seno di θ_{inc} insieme alla [retta](#) che approssima meglio i dati sperimentali, la cui pendenza vale $m \simeq 0.75$ con un errore molto piccolo e l'intercetta q è compatibile con zero.

Evidentemente, se il raggio partisse dall'acqua e si propagasse poi in aria la relazione sarebbe quella inversa, per cui, invece di avere $\sin \theta_{rifr} = 0.75 \sin \theta_{inc}$ avremmo $\sin \theta_{rifr} = 1.33 \sin \theta_{inc}$. Inoltre, se invece che propagarsi in aria, il raggio laser si propagasse in un altro materiale come il plexiglass, avremmo $\sin \theta_{rifr} = 1.12 \sin \theta_{inc}$. È chiaro che il coefficiente angolare della retta dipende dal tipo di materiale attraversato. Possiamo perciò attribuire a ciascun

materiale una caratteristica, che chiameremo **indice di rifrazione**, scrivendo la relazione che lega i seni dei due angoli come

$$n_1 \sin \theta_1 = n_2 \sin \theta_2 \quad (5.14)$$

dove abbiamo sostituito gli indici 1 e 2 ai pedici *inc* e *rifr*, rispettivamente, e n_1 e n_2 rappresentano gli indici di rifrazione (adimensionali) dei due mezzi in questione. In questa forma, la legge prende il nome di **Legge di Snell**⁵. In questo modo, nel caso del passaggio da aria ad acqua, per avere $\sin \theta_{rifr} = 0.75 \sin \theta_{inc}$, che corrisponde a $\sin \theta_2 = 0.75 \sin \theta_1$, dobbiamo assumere che $n_1/n_2 = 0.75$. Così abbiamo che $n_2/n_1 = 1/0.75 = 1.33$ e anche l'esperimento in cui la luce passa dall'acqua all'aria è verificato. Quando la luce passa dall'acqua al plexiglass avremmo

$$n_{acqua} \sin \theta_{acqua} = n_{plexiglass} \sin \theta_{plexiglass} \quad (5.15)$$

per cui $n_{plexiglass}/n_{acqua} = 1.12$. Se conveniamo di attribuire all'aria il valore $n_{aria} = 1$, allora $n_{acqua} = 1.33$ e $n_{plexiglass} = 1.49$.

5.4 Conferma della teoria corpuscolare

È compatibile questo comportamento con la teoria corpuscolare della luce? La risposta è sì e il motivo è piuttosto semplice: se la luce fosse costituita di un flusso di corpuscoli che si muovono tutti nella stessa direzione alla stessa velocità, alcuni di questi giungerebbero in prossimità della superficie di separazione tra i due mezzi prima degli altri, se l'angolo d'incidenza è diverso da zero. Se la velocità dei corpuscoli dipende dal mezzo, i primi a penetrare si muovono a velocità diversa e la direzione complessiva del flusso cambia, come si può ben vedere dal Filmato 5.10

Facendo esperimenti più raffinati si vede che l'indice di rifrazione dipende dal colore della luce. Il che

⁵Dal nome di **Willebrord Snellius**, studioso olandese vissuto tra il 1500 e il 1600.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 5.10 Un flusso di corpuscoli che si muove a velocità diversa in due mezzi diversi è **rifratto** passando da un mezzo all'altro (cambia cioè direzione): i corpuscoli che entrano nel mezzo azzurro rallentano, pur continuando a muoversi in linea retta. Si conseguenza il fronte dei corpuscoli cambia direzione [<http://youtu.be/8anT0Fs39Jo>].

non contraddice il modello, perché si può continuare a spiegare il comportamento della luce assumendo che l'indice di rifrazione dipenda dal *colore* del corpuscolo luminoso. In questo modo si spiega anche il fenomeno della **dispersione** illustrato al paragrafo precedente: se la luce bianca, che abbiamo supposto essere formata da un miscuglio di corpuscoli di *colori* diversi, incide sulla superficie di separazione tra aria e vetro, i raggi di luce giungono alla faccia opposta con angoli diversi. Se le facce di entrata e di uscita sono parallele, la rifrazione tra vetro e aria rimescola i raggi colorati in maniera da riprodurre la condizione iniziale di entrata e la luce trasmessa è bianca. Se invece le facce non sono parallele, come nel caso di un prisma, la rifrazione tra vetro e aria mantiene separati i diversi colori e si osserva quello che si chiama lo **spettro** della luce bianca.

Quando la luce giunge in prossimità della superficie di separazione tra due mezzi, una parte di essa è riflessa, mentre una parte viene rifratta, con percentuali diverse dipendenti dall'angolo d'incidenza. Potete facilmente rendervene conto provando a osservare la vostra immagine riflessa dal vetro di una finestra, la cui intensità dipende dall'angolo d'incidenza della luce. In ogni caso è sempre possibile vedere la vostra immagine al di là della finestra perché una parte della luce è rifratta. Si potrebbe interpretare il fatto assumendo che una parte del flusso di corpuscoli che giungono sulla superficie *rimbalza*,

mentre una parte è trasmessa secondo la Legge di Snell, se il mezzo è trasparente. Più l'angolo d'incidenza è radente, più la probabilità di rimbalzare è alta, come quando si tira un sasso sulla superficie del mare.

Se s'interpreta nel modo detto la luce, l'indice di rifrazione deve in qualche modo rappresentare una misura della velocità della luce che dev'essere diversa da mezzo a mezzo. Questo vuol dire che la velocità della luce non è infinita e che quindi si deve poter misurare (certo, è alta, e non sarà facile farlo). Il rapporto tra due indici di rifrazione dev'essere uguale al rapporto tra le velocità. Se quindi indichiamo con la lettera c la velocità della luce in aria abbiamo che

$$n = \frac{c}{v} \quad (5.16)$$

dove v è la velocità con la quale si propaga la luce nel mezzo di indice di rifrazione n . Sperimentalmente non si riesce a trovare alcun mezzo il cui indice di rifrazione sia minore di quello dell'aria, il che significa che la velocità della luce in aria è la più alta tra quelle conosciute e vale c .

Analizziamo meglio la Legge di Snell, scrivendo l'angolo di rifrazione di un raggio luminoso come

$$\sin \theta_{rifr} = \frac{n_{inc}}{n_{rif}} \sin \theta_{inc}. \quad (5.17)$$

Se $n_{inc} > n_{rifr}$ il coefficiente davanti a $\sin \theta_{inc}$ è maggiore di 1 e, se l'angolo d'incidenza è abbastanza grande, può accadere che $\sin \theta_{rifr} > 1$, il che è impossibile. In questo caso quel che deve accadere è che il raggio non può essere rifratto e può solo essere riflesso. Devono quindi esistere dei casi in cui, nonostante il materiale verso cui è puntato un raggio luminoso sia trasparente, il raggio non è trasmesso, ma riflesso. In effetti in certi casi è proprio quel che succede. Il fenomeno, detto **riflessione totale**, si può verificare solo quando la luce passa da un mezzo con indice di rifrazione alto a un altro con indice di rifrazione più basso. Infatti, se un raggio incide su una superficie passando da un mezzo a indice di rifrazione più alto a uno con indice di rifrazione minore, l'angolo di rifrazione è più ampio di quello

d'incidenza (il raggio si allontana dalla superficie), al contrario di quanto avviene nel caso opposto. Esiste quindi un angolo d'incidenza per il quale l'angolo di rifrazione è tale da far viaggiare l'angolo rifratto parallelamente alla superficie. Per angoli maggiori il raggio che sarebbe rifratto dovrebbe trovarsi nello stesso mezzo di partenza, ma in questo caso non potrà che formare un angolo con la normale alla superficie uguale a quello d'incidenza e quindi sarà semplicemente riflesso.

L'angolo θ_{inc} per cui si raggiunge il valore massimo di $\sin \theta_{rifr} = 1$ si chiama **angolo critico**.

5.5 Applicazioni

Quando la superficie di separazione tra due mezzi non è piana i raggi che incidono lungo una determinata direzione sono deviati in modo diverso, secondo l'angolo formato con la normale alla superficie. In particolare, se la superficie in questione è sferica, i raggi di luce paralleli all'asse ottico (che analogamente al caso degli specchi è un asse perpendicolare alla superficie) sono deviati verso quest'ultimo se incidono sulla superficie convessa, mentre si allontanano dall'asse ottico se la luce incide sulla superficie dal lato concavo. Sfruttando questo fenomeno si realizzano le **lenti**. Le lenti sono dispositivi ottici formati da un materiale trasparente con due superfici sferiche opposte. Il raggio di curvatura di ciascuna superficie ne determina le caratteristiche complessive.

Il comportamento di una lente si può studiare esattamente ricorrendo alle leggi dell'ottica geometrica sopra illustrate, ma nella maggior parte dei casi pratici si possono fare utili approssimazioni che semplificano molto la previsione del comportamento di una lente. Le approssimazioni in questione sono valide per le così dette **lenti sottili**, cioè per quelle lenti per le quali la distanza tra le due superfici opposte è molto minore del modulo di ciascuno dei due raggi di curvatura. Secondo la curvatura delle superfici le lenti possono essere biconvesse (cioè convesse su entrambi i lati), biconcave (concave su

tutti e due i lati), piano-convesso, piano-concavo o concavo-convesso.

Le lenti più comuni, per le quali è facile ricavare il comportamento, sono quelle biconvesse e biconcave. Il primo tipo devia i raggi paralleli all'asse ottico in modo da farli convergere in un punto di quest'ultimo detto **fuoco**, a distanza f dal centro della lente. Le **lenti d'ingrandimento** sono lenti di questo tipo. Disponendone l'asse ottico parallelamente alla direzione da cui provengono i raggi del Sole, questi deviano verso l'asse ottico della lente e finiscono per concentrarsi su un punto distante f dal centro della lente. Per questo motivo le lenti biconvesse sono anche dette **convergenti**. La proprietà secondo la quale i raggi paralleli all'asse ottico finiscono per convergere in un punto è quella che permette di usare una **lente d'ingrandimento** per **bruciare** un pezzo di carta usando i raggi del Sole. Poiché il Sole è molto lontano i suoi raggi arrivano sulla lente tutti parallelamente l'uno all'altro. Orientando opportunamente la lente si vede che i raggi convergono tutti in un punto, che quindi diventa luminosissimo, a distanza f dalla lente. La luce trasporta anche calore, quindi il punto in cui si concentrano i raggi, oltre a essere luminosissimo, si scalda parecchio fino a far bruciare la carta.

Perché le lenti d'ingrandimento fanno apparire più grandi gli oggetti? Per capirlo occorre tracciare il percorso di alcuni raggi di luce che partono dall'oggetto osservato e attraversano la lente. Ogni lente sottile ha due fuochi, posti simmetricamente rispetto al suo centro. Tra gli infiniti raggi che partono da un punto dell'oggetto tre in particolare sono particolarmente interessanti:

1. il raggio che viaggia parallelamente all'asse ottico: questo raggio è deviato verso il fuoco della lente, dalla parte opposta a quella in cui si trova l'oggetto;
2. il raggio che viaggia in direzione del centro della lente: questo raggio non è deviato e prosegue parallelamente a sé stesso nella regione al di là dell'oggetto;
3. il raggio che passa dal fuoco della lente che si trova dalla stessa parte dell'oggetto: questo raggio è deviato in modo da viaggiare dall'altro

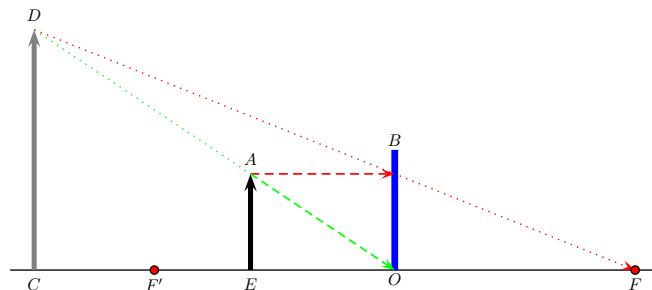


Figura 5.11 Costruzione dell'immagine prodotta da una lente convergente usata come lente d'ingrandimento. Nella figura è rappresentata solo metà della lente.

lato della lente parallelamente all'asse ottico. Come nel caso degli specchi, l'immagine si forma laddove i raggi partiti da uno stesso punto s'incontrano. Basta considerare due qualunque dei tre raggi sopra considerati per costruire l'immagine prodotta da una lente. Prendiamo in esame una lente d'ingrandimento.

Dal punto A di un oggetto, rappresentato come una freccia nella Figura 5.11, un raggio che viaggia parallelamente all'asse ottico incontra la lente nel punto B, ed è deviato nel fuoco a destra della lente. Un raggio che parte dallo stesso punto e viaggia in direzione del centro della lente invece non viene deviato. I raggi di luce che proseguono verso destra appaiono provenire entrambi da un punto in alto a sinistra, quindi chi osserva da destra vede la luce provenire da quel punto e vede l'oggetto ingrandito.

Se si osservano oggetti lontani attraverso una lente convergente, questi appaiono capovolti. Se fate l'esercizio di costruire l'immagine mettendo l'oggetto a una distanza dalla lente maggiore della distanza focale capirete facilmente il perché. Possiamo trovare alcune relazioni geometriche molto precise tra le grandezze che determinano la posizione e le dimensioni delle immagini di un oggetto visto attraverso una lente. Sempre utilizzando la Figura 5.11 vediamo che i triangoli DCF e BOF sono simili: di conseguenza, indicando con h_i l'altezza dell'immagine

$\overline{CD}$, con h_o l'altezza dell'oggetto, con p la distanza tra l'oggetto e la lente e con q quella tra l'immagine e la lente, possiamo scrivere che

$$\frac{h_i}{h_o} = \frac{q + f}{f}. \quad (5.18)$$

D'altra parte, anche i triangoli DCO e AEO sono simili per cui vale

$$\frac{h_i}{h_o} = \frac{q}{p}. \quad (5.19)$$

Combinando le due equazioni si ricava che dev'essere

$$\frac{q + f}{f} = \frac{q}{p}. \quad (5.20)$$

Quest'equazione si può riscrivere in una maniera più facile da ricordare dividendo tutto per q :

$$\frac{1}{f} + \frac{1}{q} = \frac{1}{p}. \quad (5.21)$$

Se poi facciamo la convenzione di adottare come origine degli assi la posizione della lente e di considerare positive le distanze dell'oggetto dalla lente e quelle dell'immagine se si formano dalla parte opposta della lente, vediamo che la distanza q va considerata negativa, quindi, usando per p e q i segni opportuni, l'**equazione della lente** si può riscrivere come

$$\frac{1}{f} = \frac{1}{p} + \frac{1}{q}. \quad (5.22)$$

Se l'oggetto si trova a una distanza $p > f$ dalla lente si vede subito che dev'essere

$$\frac{1}{q} = \frac{1}{f} - \frac{1}{p} \quad (5.23)$$

e quindi q è positiva e l'immagine si forma dalla parte opposta rispetto alla lente (questo accade nel caso degli obiettivi fotografici, per esempio, oppure per gli obiettivi dei proiettori). L'**ingrandimento** M si definisce come il rapporto tra le dimensioni dell'immagine e quella dell'oggetto e vale quindi:

$$M = \frac{h_i}{h_o} = \frac{q}{p}. \quad (5.24)$$

Se $M > 1$ l'immagine è più grande dell'oggetto altrimenti è più piccola. Se $M > 0$ l'immagine è dritta, altrimenti è capovolta.

Esercizio 5.2 L'equazione degli specchi

Provate a ricavare l'equazione che lega la distanza p dell'oggetto dal vertice di uno specchio sferico alla distanza q alla quale si forma l'immagine dell'oggetto e alla sua distanza focale f . Per farlo è sufficiente individuare almeno due coppie di triangoli simili disegnando i raggi che partono da uno stesso punto dell'oggetto e che viaggiano parallelamente all'asse ottico, in direzione del centro dello specchio o in direzione del suo fuoco. Scrivendo il rapporto tra h_i e h_o per due coppie di triangoli simili si ricava l'**equazione degli specchi**.

L'equazione della lente vale sia per le lenti convergenti che per quelle **divergenti**, che sono quelle biconcave. Per costruire l'immagine di una lente divergente basta seguire le stesse regole definite per le lenti convergenti, avendo cura di prolungare i raggi rifratti in modo tale che i loro prolungamenti s'incontrino in un punto, che sarà quello da cui sembrano partire i raggi luminosi e quindi quello in cui si formerà l'immagine del punto in questione. Per rendere valida l'equazione delle lenti, nel caso delle lenti divergenti basterà fare la convenzione secondo la quale il fuoco di una lente convergente f è positivo, mentre nel caso delle lenti divergenti la distanza focale si deve considerare negativa.

Agli indirizzi <http://tube.geogebra.org/student/m774117> e <http://tube.geogebra.org/student/m774113> si trovano due applicazioni **GEOGEBRA** che permettono di simulare il comportamento dei due tipi di lente.

Combinando gli effetti di due o più lenti si possono realizzare strumenti come i **telescopi** o i **microscopi**. Nei modelli più semplici, in questi strumenti ci sono solo due lenti: una che funge da **obiettivo**, da cui entra la luce che forma un'immagine a una certa distanza che dipende dalla distanza focale di questo, e una che serve da **oculare** che funziona co-

me una comune lente d'ingrandimento e che serve per consentire all'occhio di vedere l'immagine prodotta dall'obiettivo ingrandita. Poiché le immagini formate dalle lenti convergenti impiegate sono di norma rovesciate rispetto agli oggetti, quello che si vede attraverso questi strumenti appare al rovescio rispetto alla realtà. Questo è il motivo per il quale spostando il vetrino di un microscopio verso sinistra l'immagine si muove verso destra o alzando l'angolo di puntamento di un telescopio, quel che si vede nell'oculare sembra muoversi verso l'alto denunciando un movimento verso il basso del telescopio. Nei **binocoli** questo non accade perché le immagini sono ulteriormente rovesciate da una coppia di prismi montati tra la lente che costituisce l'obiettivo e quella che funge da oculare, grazie al fenomeno della riflessione totale.

Il fenomeno della riflessione totale si sfrutta non solo per la produzione dei prismi impiegati nei binocoli, ma anche nelle macchine fotografiche reflex. Nella macchine fotografiche reflex la luce proveniente dall'obiettivo è riflessa da uno specchio posto a 45° rispetto all'asse dell'obiettivo. In questo modo l'immagine è deviata verso l'alto, dove entra in un prisma, la cui faccia è perpendicolare alla direzione di propagazione della luce, che dunque non viene deviata. Viaggiando nel prisma, la luce ne incontra altre che formano con essa angoli tali da provocare riflessione totale e l'immagine è così deviata verso l'oculare. Attraverso questo il fotografo può osservare la scena inquadrata dall'obiettivo (contrariamente a quel che accade con le fotocamere il cui mirino è separato dall'obiettivo). Nel momento dello scatto lo specchio si solleva rapidamente per consentire alla luce di raggiungere il sensore posto dietro di esso. L'uso del prisma per riflettere l'immagine nel mirino consente di raddrizzarla (le immagini provenienti dall'obiettivo sono capovolte, essendo quest'ultimo un gruppo di lenti che, nell'insieme, si comportano come una lente convergente).

Un'altra applicazione del fenomeno della riflessione totale consiste nella produzione delle **fibre ottiche**. Una fibra ottica è un cilindro di materiale trasparente, detto *core*, rivestito da una pellicola di materiale (*cladding*) con indice di rifrazione più bas-

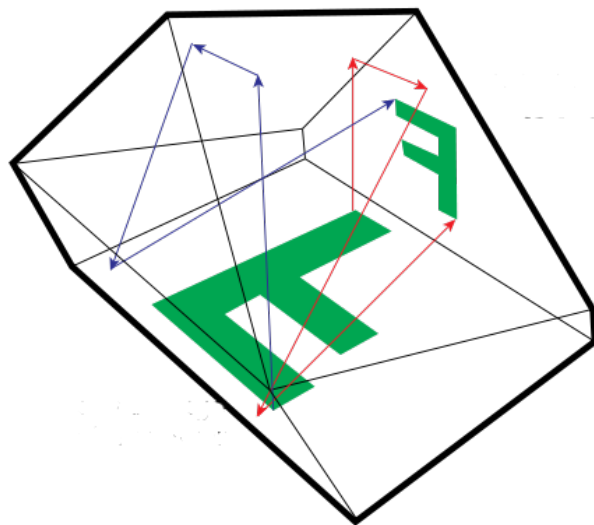


Figura 5.12 L'immagine che si forma da un obiettivo fotografico è capovolta. Per consentire al fotografo di vederla dritta, l'immagine è inviata su un prisma a sette facce in modo tale che le riflessioni totali subite dalla luce siano tali da formare l'immagine dritta in corrispondenza del mirino. Il nome **pentaprisma** dato a questo dispositivo deriva dalla forma dei primi dispositivi del genere, costituiti di prismi retti a base pentagonale (Immagine da [WIKIPEDIA](#))

so di quello del core. Quando la luce immessa nella fibra giunge alla superficie di separazione tra i due materiali, il rapporto degli indici di rifrazione è tale da favorire il verificarsi del fenomeno della riflessione totale e la luce continua a viaggiare nella fibra anche se questa forma delle curve. In questo modo è possibile condurre la luce attraverso percorsi non rettilinei, anche a grande distanza. Le fibre ottiche sono impiegate nella trasmissione di segnali per le telecomunicazioni e per Internet.

Le onde e i fenomeni ondulatori

Capita frequentemente, in fisica, di doversi occupare di fenomeni che producono l'alterazione delle caratteristiche di un mezzo determinate dal propagarsi, all'interno di quel mezzo, di un qualche tipo di sollecitazione. Quando, ad esempio, seguiamo la caduta di una moneta, di quel fenomeno c'interessa la descrizione delle grandezze fisiche che possiamo associare alla moneta che cade, che è il soggetto del nostro interesse. Della moneta possiamo tracciare la traiettoria, definire la posizione e la velocità a ogni istante di tempo. Quando la moneta raggiunge il pavimento produce un suono che, al contrario della moneta, si propaga in ogni direzione: non è qualcosa di cui possiamo tracciare la traiettoria come nel caso della moneta, ma è qualcosa di *esteso* nello spazio per il quale non sapremmo definire una procedura chiara per misurarne la *posizione*. In un certo senso potremmo dire che il suono è *ovunque* attorno all'oggetto che lo produce. Il suono, infatti, è prodotto dalla sollecitazione che la moneta, urtando il pavimento, produce in un mezzo come l'aria circostante, che si propaga nel mezzo stesso. Se togliessimo l'aria (e questo si può fare facendo il vuoto in un tubo nel quale si lascia cadere la moneta) non sentiremmo alcun rumore! In altre parole quello che *si muove*, si propaga, nel caso del suono non è qualcosa di *tangibile*, ma la sollecitazione esercitata inizialmente sul mezzo stesso. In un certo senso il mezzo attraverso il quale si propaga l'onda resta fermo (alcune sue parti si muovono, ma oscillando, per cui in media le sue parti rimangono ferme): è la perturbazione indotta nel mezzo che si propaga, si sposta, e non le parti del mezzo.

Qualcosa di analogo accade con le onde sull'acqua. Se riempiamo una bacinella d'acqua e la lasciamo ferma sul tavolo la superficie dell'acqua si di-

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 6.1 Un'onda trasversale si propaga lungo una serie di bastoncini collegati tra loro. Nel filmato si vede chiaramente come l'onda, costituita di un singolo impulso, si propaga verso sinistra, ma se si concentra l'attenzione su un singolo bastoncino si vede che esso non cambia la propria posizione, ma oscilla nella direzione verticale [http://youtu.be/0Cwtavu_E_A].

sponde a formare un piano perfettamente orizzontale e liscio. Se però infiliamo un dito nell'acqua spingiamo le particelle d'acqua verso il basso; queste, a loro volta, spingono le particelle vicine che non possono che muoversi verso l'alto e infatti così fanno. Subito dopo, però, ricadono e nel cadere spingono le particelle vicine verso il basso innescando così un processo nel quale i punti sulla superficie dell'acqua si muovono in su e in giù mentre la spinta che provoca questo movimento si propaga radialmente dal punto nel quale abbiamo infilato il dito all'inizio. La direzione di propagazione dell'**onda** che si genera è in questo caso perpendicolare¹ al moto delle particelle del mezzo attraverso il quale l'onda si propaga. In questi casi diciamo che l'onda è **trasversale** (Filmato 6.1).

Il punto di partenza di un'onda lo chiamiamo **sor-**

¹Per essere precisi il moto è un po' più complicato, ma in questa sede possiamo trascurare questo genere di dettagli.

gente dell'onda. Se la sorgente è piccola rispetto alle dimensioni dell'onda possiamo considerarla puntiforme. Quando la sorgente è puntiforme e le onde si propagano su una superficie, come nel caso in cui intingiamo un dito in una bacinella d'acqua, le onde che da questa si propagano sono di forma circolare. Quando invece le onde si propagano in un mezzo tridimensionale, nel caso di sorgenti puntiformi, le onde sono di forma sferica: quando parliamo, le onde sonore s'irradiano in tutte le direzioni rispetto alla nostra bocca e ci possono sentire anche persone che si trovano dietro, sopra o sotto di noi. Proprio questo fatto dimostra che il suono dev'essere prodotto da un'onda: se fosse prodotto da particelle che si muovono dalla sorgente all'ascoltatore dovrebbero raggiungere quest'ultimo solo nella direzione nella quale si orienta la bocca, come quando si soffia sulle candeline di una torta, che non si spengono se si trovano ai lati o dietro la bocca. Il filmato 9.1 dimostra che il suono è prodotto da un'onda **longitudinale**, nella quale cioè la sollecitazione si produce nella stessa direzione della propagazione dell'onda: una particella di aria *spinta* dalla vibrazione di un altoparlante comincia a muoversi e così facendo urta un'altra particella. La prima rimbalza all'indietro, mentre la seconda si sposta in avanti. In questo modo la seconda può trasmettere la sollecitazione a una particella vicina e così via: possiamo immaginare l'aria come un fluido che viene compresso e che torna a rarefarsi in continuazione in un determinato punto dello spazio, mentre l'onda di compressione si propaga mettendo in moto (comprimendole) le particelle di aria circostanti. Quando l'onda di compressione giunge nelle vicinanze delle nostre orecchie le molecole di aria compresse, riespandendosi, urtano il timpano e lo fanno vibrare. È questo movimento del timpano che produce nel nostro cervello la sensazione sonora.

6.1 Caratterizzazione delle onde

Se vogliamo capire di più circa il comportamento delle onde dobbiamo per prima cosa individuare al-

cune grandezze fisiche che si possono misurare nel caso delle onde. Il compito, a prima vista, non sembra facile: di onde ne esistono un'infinità e sembrano tutte diverse tra loro. Ma per fortuna la matematica ci viene in aiuto fornendoci uno strumento potentissimo: il **teorema di Fourier**². Semplificando molto, il teorema di Fourier stabilisce che ogni funzione **periodica** (con certe caratteristiche come il fatto di essere **continua**, che però è una caratteristica di cui godono praticamente tutte le funzioni che interessano a un fisico), si può sempre scrivere come una somma, al limite infinita, di funzioni sinusoidali. In formule:

$$f(x) = \sum_{i=0}^{\infty} A_n \cos \frac{x}{L_n} + B_n \sin \frac{x}{L_n}, \quad (6.1)$$

dove A_n , B_n e L_n sono opportuni coefficienti numerici che il teorema consente di calcolare. Naturalmente la stessa formula si può scrivere in molti modi diversi e nei libri di matematica appare spesso in un'altra forma, ma voi avrete notato che l'argomento delle funzioni seno e coseno è stato scritto come x/L_n per evidenziarne la natura adimensionale: le dimensioni fisiche di x devono evidentemente essere le stesse di L_n . Le funzioni periodiche sono quelle che si ripetono uguali a sé stesse a intervalli regolari. Per rendervi conto meglio di come funzioni il teorema di Fourier potete guardare i grafici animati che si trovano alla pagina di Wikipedia che corrisponde alla voce **Fourier Series**, che è fatta piuttosto bene. In particolare a quella pagina si trova la Fig. 6.2.

Potete anche accedere a <http://tube.geogebra.org/student/m811701> per sperimentare da soli la decomposizione di una funzione periodica in **armoniche**, come sono dette le componenti della somma di Fourier.

Le onde cui siamo interessati noi sono sempre rappresentabili come funzioni periodiche. Anche qualora l'onda sia costituita di un singolo impulso (vedi Filmato 6.1), come quando si scuote l'estremità di

²Dal nome di **Jean Baptiste Fourier** che lo formulò.

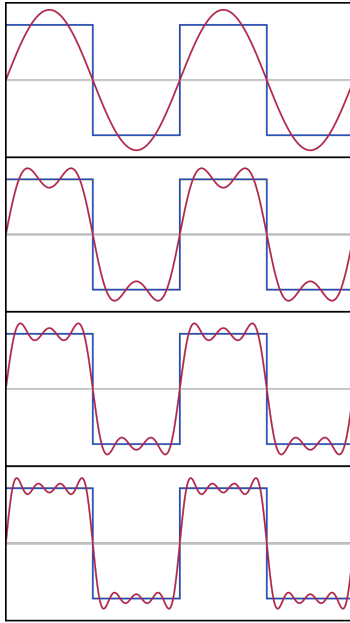


Figura 6.2 Un'onda quadra, che si presenta come una successione regolare di impulsi di forma rettangolare si può sempre scrivere come una somma di tante funzioni sinusoidali. In alto si vede l'onda quadra, in blu, sovrapposta a una singola senoide: le due onde si somigliano, ma non troppo. Se però alla senoide se ne somma un'altra con un diverso periodo si trova la figura immediatamente sotto, che decisamente somiglia di più all'onda in blu. Continuando a sommare sinusoidi si ottiene una funzione che somiglia sempre di più a quella originale (La figura è stata prodotta da Jim Belk).

una corda³, si può pensare come rappresentata da una funzione periodica il cui periodo, cioè l'intervallo di tempo dopo il quale un secondo impulso compare laddove è transitato il primo, è molto lungo (al limite infinito). Questo ci permette di studiare, dal punto di vista formale, solo le onde sinusoidali:

³L'impulso si muove lungo la corda, ma i punti della corda non si spostano da dove sono: vanno semplicemente su, poi giù.

quelle cioè che si possono rappresentare come funzioni trigonometriche seno o coseno. Tanto, qualunque sia la forma dell'onda si potrà sempre scrivere come somma di onde trigonometriche e dunque, nel complesso, l'onda avrà le stesse proprietà delle onde sinusoidali. Ci possiamo dunque limitare a considerare solo onde sinusoidali: si tratta di una semplificazione non da poco! Nel caso delle onde sulla superficie dell'acqua, per esempio, possiamo pensare di limitarci a descriverle matematicamente con una funzione del tipo

$$y = A \sin \theta, \quad (6.2)$$

dove y rappresenta una grandezza fisica il cui valore cambia al variare di qualcosa che si può rappresentare come un angolo. Per esempio, nel caso delle onde sull'acqua, y rappresenta l'altezza del pelo dell'acqua rispetto alla condizione di equilibrio e può essere sia positiva (sulle **creste** dell'onda) che negativa (nei **ventri**). Se l'onda descrive un suono, y rappresenta la variazione di pressione del fluido nel quale il suono si propaga. La pressione in un determinato punto può essere maggiore di quella media (e in questo caso $y > 0$) o minore ($y < 0$).

Come *argomento* del seno non possiamo che usare un angolo θ che di fatto sarà un numero adimensionale compreso tra $-\infty$ e $+\infty$ che rappresenta una grandezza fisica da cui l'onda dipende: una distanza o un intervallo di tempo. L'ampiezza dell'onda y a un certo istante di tempo t , infatti, dipende dalla posizione x in cui si osserva il mezzo sollecitato. Prendiamo per esempio il caso di un'onda che si propaga lungo una corda. Se copriamo tutta la corda con uno schermo con una finestra sufficientemente sottile e osserviamo i punti della corda attraverso la finestra, li vedremo andare su e giù al variare del tempo. In questo caso y è una funzione del tempo $y = y(t)$ e x è fissata alla posizione della finestra. Se invece fissiamo il tempo, facendo una foto della corda dopo aver rimosso lo schermo, vedremo che l'altezza dei punti della corda y è funzione della distanza da uno dei capi della corda $y = y(x)$, mentre il tempo t è fissato all'istante in cui la foto è stata fatta.

Di un'onda (sinusoidale) possiamo misurare diverse caratteristiche. Una di queste è l'**ampiezza** A che rappresenta la deviazione massima dei punti del mezzo sollecitato dall'onda dalla posizione di equilibrio (nel caso delle onde sull'acqua l'altezza massima che raggiunge l'onda rispetto al caso in cui l'acqua è immobile). Se concentriamo la nostra attenzione su un particolare istante di tempo (se quindi facciamo una fotografia all'acqua mentre si propaga l'onda) possiamo definire chiaramente un'altra caratteristica che è la **lunghezza d'onda** λ , che rappresenta la distanza che separa due punti della sinusoide che abbiano le stesse caratteristiche di ampiezza e di pendenza: la lunghezza d'onda è la distanza, in metri, tra due massimi consecutivi, che è identica alla distanza tra due minimi o tra due punti di ampiezza nulla nei quali però l'onda stia crescendo o decrescendo in entrambi (due punti consecutivi nei quali $y = 0$ distano metà di una lunghezza d'onda). Se viceversa ci concentriamo su un particolare punto dell'acqua e misuriamo il tempo che l'acqua impiega a raggiungere la massima altezza, tornare giù, raggiungere la minima altezza e tornare al punto di partenza, troviamo un tempo che chiamiamo **periodo** e indichiamo con T .

In definitiva l'ampiezza rappresenta la variazione massima (in valore assoluto) delle caratteristiche variabili del mezzo sollecitato (l'altezza dell'acqua nel caso delle onde sull'acqua, la densità dell'aria nel caso delle onde sonore, etc.), la lunghezza d'onda la distanza tra due punti consecutivi dell'onda che abbiano le stesse caratteristiche di ampiezza e di pendenza nello stesso istante e il periodo l'intervallo di tempo necessario affinché l'onda si ripeta uguale a sé stessa nello stesso punto.

Concentriamo la nostra attenzione su un massimo di un'onda sinusoidale e seguiamone l'evoluzione temporale: se al tempo $t = 0$ il massimo si trova in $x = 0$, al tempo $t = T$ si troverà nel punto di coordinate $x = \lambda$. Al tempo $t = T/2$ avrà raggiunto il punto di coordinate $x = \lambda/2$ ed è facile convincersi che al tempo $t = T/\alpha$ il massimo che stiamo seguendo si sarà spostato di $x = \lambda/\alpha$ per cui il rapporto x/t vale

$$\frac{x}{t} = \frac{\lambda}{\alpha T} = \frac{\lambda}{T} = v \quad (6.3)$$

ed è costante. Questo rapporto lo chiamiamo **velocità** dell'onda. La velocità di un'onda ci dice quanto rapidamente si sposta un massimo (o un punto qualsiasi con le stesse caratteristiche di ampiezza e pendenza): maggiore è la velocità minore è il tempo impiegato dall'onda a spostarsi di una stessa distanza. Dal momento che le dimensioni fisiche di λ sono quelle di una distanza e quelle di T sono quelle di un tempo, le dimensioni fisiche di v sono quelle di una distanza divisa per un tempo e si misura quindi in m/s nel SI.

È anche utile definire l'inverso di T , $\nu = 1/T$, come la **frequenza** dell'onda. La frequenza, che evidentemente ha le dimensioni di un tempo alla meno uno, si misura in s^{-1} o in Hertz (Hz: 1 Hz = 1 s^{-1}).

L'angolo θ nell'equazione (6.2) è un angolo che deve dipendere dalla posizione x e dall'istante di tempo t in cui si osserva l'onda. Se guardiamo l'onda in un preciso istante di tempo (ad esempio se facciamo una fotografia dell'onda) vediamo che, se scegliamo il punto $x = 0$ nel punto in cui l'ampiezza dell'onda è la massima possibile A , allora l'onda dovrà assumere nuovamente il valore massimo per $\theta = 2\pi$ e quindi la distanza x alla quale si trova⁴ sarà tale da rendere quest'angolo pari all'angolo giro. Ma la distanza x alla quale l'onda assume di nuovo il valore massimo è λ perciò dev'essere

$$\theta = x \frac{2\pi}{\lambda} \quad (6.4)$$

in modo tale che per $x = 0$, $\theta = 0$ e che per $x = \lambda$, $\theta = 2\pi$. In questo modo per $x = \lambda/2$, $\theta = \pi$ e così via. L'equazione dell'onda dunque si scrive come

$$y = A \sin \frac{2\pi}{\lambda} x. \quad (6.5)$$

D'altra parte, se invece concentriamo la nostra attenzione su un singolo punto del mezzo nel quale

⁴Con quest'espressione non troppo precisa intendiamo la posizione alla quale si trova un punto dell'onda scelto come riferimento: può trattarsi di una cresta, di un ventre o di un altro punto qualunque.

si sta propagando l'onda, lo vedremo oscillare: nel caso delle onde sull'acqua, per esempio, lo vedremo andare su e giù. Se scegliamo l'istante $t = 0$ quando la grandezza fisica che descrive le caratteristiche del punto del mezzo assume il valore $y = A$, che è il massimo possibile (per esempio l'altezza rispetto al fondo di una vasca nel caso delle onde sull'acqua), il punto tornerà ad assumere quelle stesse condizioni dopo un tempo $t = T$ il che corrisponde a un angolo $\theta = 2\pi$. Pertanto possiamo scrivere che

$$\theta = t \frac{2\pi}{T} \quad (6.6)$$

in modo tale che per $t = 0$, $\theta = 0$ e per $t = T$, $\theta = 2\pi$. In questo caso l'onda si scrive come

$$y = A \sin \frac{2\pi}{T} t. \quad (6.7)$$

Sapendo che $T = 1/\nu$ possiamo riscrivere l'equazione dell'onda nella forma

$$y = A \sin (2\pi \nu t) \quad (6.8)$$

e definendo la **pulsazione** $\omega = 2\pi\nu$, come

$$y = A \sin \omega t. \quad (6.9)$$

Quando $x \neq 0$ evidentemente sarà

$$y = A \sin \left(\frac{2\pi}{\lambda} x \pm \frac{2\pi}{T} t \right) \quad (6.10)$$

dove il segno $+$ o $-$ dipende dal verso nel quale si sta propagando l'onda. Se l'onda si propaga nel verso delle x positive, allora dopo un tempo $t = T/2$ il massimo deve essersi spostato in $x = \lambda/2$, quindi l'argomento del seno deve valere zero quando $t = T/2$ e $x = \lambda/2$. Deve quindi valere la relazione

$$\left(\frac{4\pi}{\lambda} \lambda \pm \frac{4\pi}{T} T \right) = 0 \quad (6.11)$$

e questo si può verificare solo se il segno è quello negativo. Viceversa, se il segno è positivo l'onda si sta propagando verso i punti con x più piccole. Ricordando che $v = \lambda/T$ (quindi $T = \lambda/v$) possiamo riscrivere l'onda sostituendo a T il valore $T = \lambda/v$:

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 6.3 Nel filmato si vede un'onda che si propaga da destra a sinistra e che, giunta alla fine del mezzo nel quale si sta propagando, è riflessa all'indietro [<http://youtu.be/NcNR7XsqPc>].

$$y = A \sin \left(\frac{2\pi}{\lambda} x \pm \frac{2\pi}{T} t \right) = A \sin \left(\frac{2\pi}{\lambda} (x \pm vt) \right). \quad (6.12)$$

Il **fronte d'onda** è il luogo dei punti del mezzo in cui il valore di y (che, ricordiamo, è la grandezza fisica che caratterizza l'onda) è lo stesso. Il fronte d'onda avanza nel tempo spostandosi a x maggiori o minori, secondo il segno davanti a v , quindi la direzione dello spostamento del fronte è sempre perpendicolare al fronte stesso.

6.2 Riflessione e rifrazione delle onde

Anche le onde, come i corpuscoli, sono soggette ai fenomeni della **riflessione** e della **rifrazione**. Quando un'onda incontra un ostacolo sufficientemente grande e con le opportune proprietà può essere **riflessa**. La perturbazione che si propaga, giunta in prossimità dell'ostacolo, cambia direzione o verso. Se l'onda arriva perpendicolarmente sull'ostacolo viene riflessa all'indietro. Se invece arriva sull'ostacolo formando un angolo θ_i con la direzione normale alla superficie d'incidenza, l'onda è riflessa in modo tale che l'angolo di riflessione θ_r sia uguale all'angolo d'incidenza θ_i : $\theta_i = \theta_r$.

Se invece l'onda incide sulla superficie di separazione tra due mezzi nei quali la velocità di propagazione delle onde non è la stessa, subisce il fenomeno della rifrazione che consiste nel cambio di direzione dell'onda. La direzione dell'onda dopo aver attraversato la superficie tende ad avvicinarsi alla normale

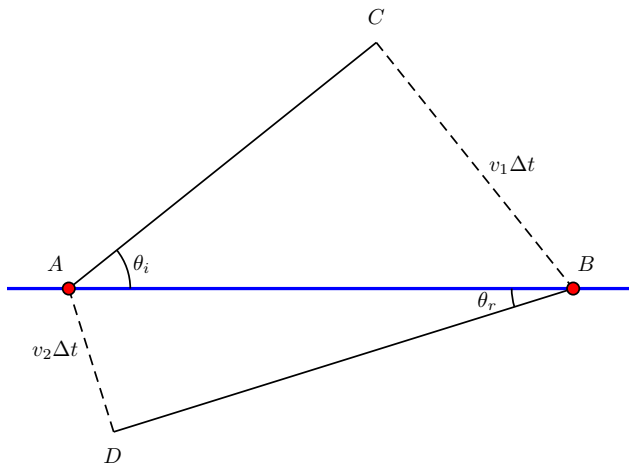


Figura 6.4 Il fronte d'onda $\overline{AC}$ giunge in prossimità della superficie di separazione tra due mezzi (uno dei quali rappresentato in azzurro). Il fronte d'onda avanza nel mezzo 2 con velocità diversa e nello stesso tempo nel quale il fronte percorre il tratto $\overline{CB}$ nel mezzo 1, percorre il tratto $\overline{AD}$ nel mezzo 2.

alla superficie se nel secondo mezzo la velocità di propagazione è minore che nel primo, e tende ad allontanarsene se invece è maggiore.

Questo è uno dei motivi per i quali le onde che si vedono arrivare su una spiaggia arrivano sempre da una direzione che è grosso modo perpendicolare alla linea di costa. In alto mare le onde possono avere una direzione qualunque, quindi potrebbero anche viaggiare quasi parallelamente alla costa, ma avvicinandosi a questa la loro velocità diminuisce⁵ e così la direzione di propagazione tende ad avvicinarsi sempre di più a quella perpendicolare alla spiaggia.

Il fenomeno è piuttosto semplice da spiegare: quando il fronte d'onda arriva con un angolo θ_i rispetto alla direzione normale alla superficie di propagazione, la parte di questo più avanzata penetra

⁵La velocità delle onde del mare dipende dalla profondità del fondale.

per prima nel secondo mezzo e rallenta, restando indietro rispetto alla porzione di fronte d'onda che ancora non è arrivata alla superficie. Man mano che il fronte d'onda prosegue, il rallentamento interessa una porzione sempre più grande del fronte. Facendo riferimento alla Fig. 6.4 si vede che al tempo $t = 0$ il fronte d'onda, rappresentato dal segmento $\overline{AC}$, giunge alla superficie di separazione; dopo un tempo Δt la parte che ha già attraversato la superficie è avanzata di $\Delta x_2 = v_2 \Delta t$, mentre quella che a $t = 0$ era ancora nel mezzo 1 ha percorso un tratto lungo $\Delta x_1 = v_1 \Delta t$ più lungo⁶. I triangoli ACB e ADB sono rettangoli, quindi possiamo scrivere che

$$\overline{AB} \sin \theta_i = v_1 \Delta t \quad (6.13)$$

e

$$\overline{AB} \sin \theta_r = v_2 \Delta t. \quad (6.14)$$

Dividendo membro a membro abbiamo che

$$\frac{\sin \theta_i}{\sin \theta_r} = \frac{v_1}{v_2}. \quad (6.15)$$

Se c è la velocità dell'onda in un mezzo di riferimento possiamo sempre esprimere v_1 e v_2 come la velocità c divisa per un numero n adimensionale che dipende dal mezzo nel quale si sta propagando l'onda: $v_1 = c/n_i$ e $v_2 = c/n_r$. Abbiamo perciò che

$$\frac{\sin \theta_i}{\sin \theta_r} = \frac{c}{n_i} \frac{n_r}{c} = \frac{n_r}{n_i}. \quad (6.16)$$

che si può riscrivere come

$$n_i \sin \theta_i = n_r \sin \theta_r \quad (6.17)$$

che mnemonicamente è più semplice da ricordare. Al numero n si dà il nome di **indice di rifrazione**.

⁶Per eseguire il disegno bisogna considerare che il fronte d'onda, per definizione, si propaga perpendicolarmente a sé stesso, quindi il tratto AD è perpendicolare alla porzione di fronte nel mezzo in basso DB , mentre il tratto CB è perpendicolare alla porzione di fronte nel mezzo in alto AC .

PREREQUISITI: Onde

Quando un mezzo è sollecitato in più punti, in ciascuno di essi si genera un'onda che si propaga. In uno stesso punto del mezzo si possono dunque produrre effetti dovuti a più onde: effetti che si sommano dando luogo al fenomeno dell'**interferenza** che consiste proprio nella somma (algebrica) degli effetti provocati da più onde in uno stesso punto.

Sollecitando la superficie di un liquido (per esempio lanciando un sasso in uno stagno), dal punto in cui avviene la sollecitazione si propagano onde circolari. Se i punti sollecitati sono due (due sassi lanciati in due punti diversi dello stagno) le onde si propagano da due punti diversi e prima o poi i fronti d'onda circolari s'intersecano. Nei punti in cui giungono contemporaneamente le creste di due diverse onde, gli effetti prodotti da ciascuna si sommano e si dice che si ha **interferenza costruttiva**. Se in uno stesso punto giungono contemporaneamente due onde ciascuna delle quali ha ampiezza nulla al momento dell'incontro, gli effetti provocati sul mezzo saranno ugualmente nulli. Può anche succedere che due onde arrivino nello stesso punto dello spazio quando si trovano l'una nella condizione di ampiezza massima e l'altra di ampiezza minima: in questo caso gli effetti dell'una annullano (almeno parzialmente se le onde non sono uguali) gli effetti dell'altra e si ha **interferenza distruttiva**. In generale possiamo dire che si ha interferenza ogni qual volta due o più onde s'incontrano nello stesso punto dello spazio. Nel medesimo istante, in quel punto il mezzo appare sollecitato da tutte le onde che vi giungono e dunque la grandezza fisica y perturbata dall'onda (l'altezza del pelo dell'acqua, la pressione locale del

fluido, etc.) risulterà avere un valore che si ottiene sommando tutti gli effetti prodotti da tutte le onde:

$$y = \sum_i A_i \sin \theta_i. \quad (7.1)$$

Ciascun θ_i si può scrivere come

$$\theta_i = \frac{2\pi}{\lambda_i} (x_i \pm v_i t_i). \quad (7.2)$$

Nell'espressione, l'ampiezza A_i di ciascuna onda è indipendente dall'altra, così come la sua lunghezza d'onda λ_i . La velocità v_i con la quale si propaga l'onda nel mezzo, in linea di principio, può essere diversa per ciascuna onda. Sperimentalmente si vede che, per quasi tutte le onde, la velocità di propagazione dipende più che altro dal mezzo nel quale si propagano e solo debolmente dalle caratteristiche dell'onda. Per esempio, la velocità con la quale si propagano le onde sull'acqua o le onde sonore risulta praticamente indipendente dalla lunghezza d'onda (nel caso del suono in aria è di circa 330 m/s), anche se una debole dipendenza c'è. In ogni caso, fino a quando non abbiamo problemi con la matematica, ci converrà assumere che la velocità di propagazione dell'onda potrebbe dipendere dall'onda stessa.

Il valore di x_i rappresenta la distanza tra la **sorgente** dell'onda i e il punto nel quale si osserva l'onda e t_i il tempo trascorso dall'istante in cui questa è partita dalla sorgente. Se vogliamo sommare tutte le onde insieme dobbiamo farlo rispetto a un unico sistema di riferimento che ha l'origine in un punto che potrebbe coincidere con una delle sorgenti, ma non con le altre. Perciò, se misuriamo la distanza percorsa dall'onda a partire da un punto arbitrario O , le distanze x_i percorse da ciascuna onda si scriveranno come

$$x_i = x + \delta_i \quad (7.3)$$

e contestualmente

$$t_i = t + \tau_i \quad (7.4)$$

con δ_i e τ_i costanti. La conseguenza è che ciascuna onda si scrive come

$$y_i = A_i \sin \left(\frac{2\pi}{\lambda_i} ((x + \delta_i) \pm v_i (t + \tau_i)) \right). \quad (7.5)$$

Raccogliendo tutte le costanti si trova

$$y_i = A_i \sin \left(\frac{2\pi}{\lambda_i} ((x \pm v_i t) + \delta_i \pm v_i \tau_i) \right) \quad (7.6)$$

e osservando che

$$\frac{2\pi}{\lambda_i} (\delta_i \pm v_i \tau_i) = \phi_i \quad (7.7)$$

è costante, possiamo riscrivere

$$y_i = A_i \sin \left(\frac{2\pi}{\lambda_i} (x \pm v_i t) + \phi_i \right) \quad (7.8)$$

La grandezza ϕ_i , che è adimensionale, prende il nome di **fase** dell'onda e dipende dalla scelta che si fa del sistema di riferimento per misurare le distanze e i tempi. La somma di tutte queste onde si scrive dunque

$$y = \sum_i y_i = \sum_i A_i \sin \left(\frac{2\pi}{\lambda_i} (x \pm v_i t) + \phi_i \right). \quad (7.9)$$

La somma di tante funzioni sinusoidali può dar luogo, in linea di principio, a una funzione che può assumere un valore massimo pari a $A_{max} = \sum_i A_i$ e un minimo di $A_{min} = -A_{max}$. Naturalmente, data la presenza della fase ϕ_i , potrebbe anche accadere che $A = \sum A_i = 0$, così come è possibile che l'ampiezza assuma uno qualunque dei valori compresi tra $-A_{max}$ e $+A_{max}$.

All'indirizzo <http://tube.geogebra.org/student/m811957> hai la possibilità di vedere cosa succede alla somma di due onde sinusoidali quando la fase relativa ϕ cambia.

7.1 Casi particolari

La somma di molte onde può essere complicata da trattare dal punto di vista matematico, perciò, per capire gli effetti provocati dall'interferenza, ci limitiamo a considerare quella prodotta da due sole onde con la stessa ampiezza. I fenomeni descritti in questo paragrafo si possono sperimentare all'indirizzo <http://tube.geogebra.org/student/m812001>. Considerando due onde con la stessa ampiezza, ma con frequenza diversa, abbiamo che, in uno stesso punto x e allo stesso istante t l'onda risultante è

$$y = A \left(\sin \left(\frac{2\pi}{\lambda_1} (x \pm v_1 t) \right) + \sin \left(\frac{2\pi}{\lambda_2} (x \pm v_2 t) + \phi \right) \right). \quad (7.10)$$

La fase ϕ l'abbiamo messa solo in uno degli addendi perché possiamo sempre scegliere il sistema di riferimento in modo tale che $\phi_1 = 0$ e $\phi_2 = \phi$. In questo caso ϕ si dice **fase relativa**. Introducendo il **numero d'onda** $k = 2\pi/\lambda$ e ricordando che $\omega = 2\pi/T$ e che $v = \lambda/T$ l'equazione dell'onda si scrive

$$y = A (\sin (k_1 x \pm \omega_1 t) + \sin (k_2 x \pm \omega_2 t) + \phi). \quad (7.11)$$

Può sembrare complicato ricordare le espressioni adoperate, ma in realtà è molto semplice: basta ricordare che l'obiettivo è quello di scrivere la funzione che rappresenta l'onda nella maniera più semplice possibile che è $y = A \sin (kx \pm \omega t)$. A questo punto il numero k dev'essere qualcosa che ha le dimensioni di una lunghezza alla meno uno e tale da formare un angolo quando è moltiplicato per x . L'unica grandezza che caratterizza l'onda che ha le dimensioni di una lunghezza è λ perciò $k \propto 1/\lambda$. Dovendo kx rappresentare un angolo è chiaro che ci vuole un fattore 2π a numeratore, perciò

$$k = \frac{2\pi}{\lambda} \quad (7.12)$$

che si misura in m^{-1} nel SI. Allo stesso modo ωt dev'essere un angolo, quindi ω dev'essere un tempo alla meno uno e non può che essere

$$\omega = \frac{2\pi}{T} = 2\pi\nu. \quad (7.13)$$

Utilizzando le formule di **prostaferesi** riscriviamo la funzione d'onda come

$$y = 2A \sin \frac{(k_1 + k_2)x \pm (\omega_1 + \omega_2)t + \phi}{2} \cos \frac{(k_1 - k_2)x \pm (\omega_1 - \omega_2)t - \phi}{2}. \quad (7.14)$$

Possiamo a questo punto misurare gli effetti dell'onda nello stesso punto x e vedere come varia nel tempo. Scegliendo l'origine degli assi nel punto in cui osserviamo l'onda abbiamo $x = 0$ e sostituendo nella funzione d'onda avremo la descrizione di ciò che vedremo in quel punto al variare del tempo: la grandezza fisica y che oscilla come

$$y = 2A \sin \left(\frac{\omega_1 + \omega_2}{2}t + \frac{\phi}{2} \right) \cos \left(\frac{\omega_1 - \omega_2}{2}t - \frac{\phi}{2} \right). \quad (7.15)$$

Nel caso in cui le due onde che interferiscono siano in fase ($\phi = 0$) e abbiano la stessa frequenza per cui $\omega = \omega_1 = \omega_2$, abbiamo che l'argomento del coseno vale 1, mentre quello del seno vale ω e quindi

$$y = 2A \sin \omega t \quad (7.16)$$

risulta essere un'onda che ha le stesse caratteristiche di quelle incidenti, ma ampiezza doppia. Se la fase è nulla, ma le due onde non hanno la stessa frequenza abbiamo che

$$y = 2A \sin \left(\frac{\omega_1 + \omega_2}{2}t \right) \cos \left(\frac{\omega_1 - \omega_2}{2}t \right). \quad (7.17)$$

Possiamo pensare al prodotto

$$A(t) = 2A \cos \left(\frac{\omega_1 - \omega_2}{2}t \right) \quad (7.18)$$

come a un'**ampiezza variabile** dell'onda risultante, che sarà descritta come

$$y = A(t) \sin \left(\frac{\omega_1 + \omega_2}{2}t \right). \quad (7.19)$$

In questo caso l'ampiezza dell'onda non è costante, ma varia nel tempo e varia sinusoidalmente con una frequenza che è la differenza delle frequenze delle due onde. L'ampiezza passa da un valore minimo a un valore massimo e in certi casi è nulla. In particolare se $\omega_1 \simeq \omega_2 \simeq \omega$ la semisomma di queste due frequenze è pari alla media delle frequenze (che sarà simile alle frequenze originarie), mentre la semidifferenza sarà un numero molto piccolo. Se la frequenza con la quale cambia l'ampiezza è piccola, la modulazione dell'ampiezza dell'onda varia lentamente nel tempo e si parla di **battimenti**.

I battimenti si possono apprezzare facilmente con le onde sonore: pizzicando la stessa corda di due chitarre ben accordate, il suono prodotto da entrambe ha la stessa frequenza e nello stesso punto si può ascoltare un suono d'intensità più alta rispetto a quello prodotto da una sola delle due. In questo caso la fase relativa è abbastanza irrilevante perché comunque dà origine a una fase complessiva costante che si traduce in una traslazione dell'origine dei tempi. Se la nota prodotta da una delle due corde è alterata agendo sulla tastiera, in modo tale che la frequenza corrispondente non sia troppo lontana da quella originale, il suono che si percepisce è praticamente lo stesso di prima, ma si ode una modulazione regolare dell'intensità del suono che appare ora più intenso ora meno intenso a intervalli regolari (determinati dalla differenza $\omega_1 - \omega_2$).

Un altro caso interessante d'interferenza si ha quando un'onda di equazione

$$y = A \sin \frac{2\pi}{\lambda}x \quad (7.20)$$

interferisce con un'altra onda con la stessa lunghezza d'onda partita da un punto diverso $x' = x + \delta$, che ha dunque equazione

$$y = A \sin \frac{2\pi}{\lambda}(x - \delta). \quad (7.21)$$

Quando le due onde si sommano si ottiene un'onda di equazione

Accordare un pianoforte

Il fenomeno dei battimenti è sfruttato dagli accordatori di strumenti musicali per regolare la tensione delle corde degli strumenti che corrispondono all'emissione della nota LA. Questa è la nota prodotta dai **diapason** quando sono percossi. Se insieme al diapason si fa vibrare la corda corrispondente, in caso di perfetta accordatura il suono che si percepisce è piatto e d'intensità costante (con una leggera diminuzione col tempo dovuta alle forze dissipative). Quando invece la nota prodotta dalla corda è vicina, ma non identica, a quella prodotta dal diapason si ode la tipica modulazione del volume dei battimenti. La percezione di questa modulazione induce l'accordatore ad aumentare o a diminuire la tensione della corda.

$$y = 2A \sin \left(\frac{2\pi}{\lambda} \left(x - \frac{\delta}{2} \right) \right) \cos \left(\frac{2\pi}{\lambda} \frac{\delta}{2} \right). \quad (7.22)$$

Ora, se l'argomento del coseno vale 2π o un suo multiplo, il coseno vale 1. Questa circostanza si verifica quando

$$\frac{2\pi}{\lambda} \frac{\delta}{2} = 2m\pi \quad (7.23)$$

dove m è un qualunque numero intero $m = 0, 1, 2, \dots$. Questo significa che quando

$$\delta = 2m\lambda, \quad (7.24)$$

e cioè quando le due onde sono sfasate di un **numero pari di lunghezze d'onda**, si ha la massima interferenza costruttiva. Del resto, anche quando l'argomento del coseno vale π (e quindi il coseno vale -1) si ha interferenza costruttiva, perché quel che succede è che l'onda risultante assume l'ampiezza massima con il segno cambiato. Questo avviene quando

$$\frac{2\pi}{\lambda} \frac{\delta}{2} = m\pi \quad (7.25)$$

e quindi per

$$\delta = m\lambda. \quad (7.26)$$

L'interferenza costruttiva quindi si ha quando le due onde sono sfasate di un **qualunque numero intero di lunghezze d'onda**. Al contrario, l'interferenza è totalmente distruttiva quando l'argomento del coseno vale $\frac{\pi}{2}, 3\frac{\pi}{2}, 5\frac{\pi}{2}$ e così via, cioè quando l'argomento del coseno è un multiplo dispari di $\frac{\pi}{2}$. Questo si verifica quando

$$\frac{2\pi}{\lambda} \frac{\delta}{2} = (2m+1) \frac{\pi}{2} \quad (7.27)$$

con m intero, cioè quando

$$\delta = (2m+1) \frac{\lambda}{2}. \quad (7.28)$$

Possiamo dunque dire che l'interferenza sarà completamente distruttiva quando le due onde sono sfasate di un **numero dispari di mezze lunghezze d'onda**.

7.2 Onde stazionarie

Quando un'onda si propaga in un mezzo e giunge in un punto nel quale il mezzo è vincolato a star fermo, deve avere, in quel punto, ampiezza nulla. Nel caso delle onde che si propagano lungo una corda, ad esempio, può succedere che la corda sia legata in uno o più punti a un supporto: in questo caso, quando si sposta un punto della corda dalla sua posizione di equilibrio e lo si lascia andare, il moto dei vari punti della corda assume carattere ondulatorio, come sempre. Le onde che si formano devono necessariamente avere ampiezza nulla nei punti in cui la corda è vincolata. La corda di un violino o di una chitarra, per esempio, da un lato è vincolata al *ponticello*, dall'altra al *piolo* (detto anche *bischero*) che è il piolo attorno al quale si avvolge e che serve per tenderla. Quando la si pizzica, si spostano alcuni suoi punti dalla posizione di equilibrio (il che è possibile grazie alle sue proprietà elastiche); in seguito al rilascio, le forze elastiche riportano i punti spostati nella posizione iniziale, ma le forze che

tengono insieme i punti della corda producono a loro volta spostamenti sui punti adiacenti che danno origine al propagarsi di due onde nei due versi della corda: una verso il ponticello e l'altra verso il pirolo. La somma delle due onde costituisce a sua volta un'onda che tuttavia non può essere qualunque: nel punto in cui la corda è assicurata al ponticello e nel punto in cui è avvolta attorno al pirolo l'ampiezza dell'onda dev'essere nulla, deve cioè essere che

$$y(x, t) = A \sin(kx + \omega t) + A \sin(kx - \omega t) = 0 \quad (7.29)$$

quando x corrisponde alla coordinata del ponticello (che possiamo sempre scegliere coincidere con $x = 0$) e quando x coincide con la coordinata del pirolo (per cui, se abbiamo scelto $x = 0$ coincidente con il ponticello, abbiamo $x = L$ dove L è la lunghezza effettiva della parte di corda che può oscillare). Questo significa che, sempre usando le formule di **prostaferesi** e il fatto che $\cos(-\theta) = -\cos \theta$, si deve avere

$$y(x, t) = -2A \sin(kL) \cos(\omega t) = 0 \quad (7.30)$$

per qualunque valore di t . Questo può accadere soltanto se

$$kL = m\pi \quad (7.31)$$

con m pari a un qualunque numero intero $m = 0, 1, 2, \dots$. Poiché $k = 2\pi\lambda^{-1}$ (per ricordarlo basta osservare che le dimensioni di k devono essere quelle di una lunghezza alla meno uno e che il prodotto kx dev'essere un angolo) abbiamo che

$$\frac{2\pi}{\lambda} = m \frac{\pi}{L} \quad (7.32)$$

cioè che

$$\frac{\lambda}{2} = \frac{L}{m} \quad (7.33)$$

e in definitiva dev'essere $\lambda = 2L/m$ con m intero. La lunghezza d'onda delle onde che si propagano

lungo una corda fissata in due punti non può essere qualunque: solo certe lunghezze d'onda sono permesse, la più grande delle quali è pari a $\lambda_{MAX} = 2L$. Corrispondentemente la frequenza dell'onda può assumere solo valori discreti e la minima frequenza permessa per le onde di questo tipo vale

$$\nu_{min} = \frac{c}{\lambda_{MAX}} = \frac{c}{2L}, \quad (7.34)$$

dove c è la velocità di propagazione dell'onda lungo la corda. La frequenza ν_{min} si chiama **fondamentale** e l'onda associata prende il nome di **armonica fondamentale**. Tutte le altre frequenze sono multiple della fondamentale essendo

$$\nu = \frac{c}{\lambda} = \frac{c}{2L}m = m\nu_{min} \quad (7.35)$$

e prendono il nome di **armoniche di ordine m**. Questo genere di onde si chiama **onda stazionaria** perché in essa alcuni punti del mezzo nel quale si propaga l'onda, che si chiamano **nodi**, sono sempre fermi.

Un secondo tipo di onde stazionarie si verifica quando il numero di punti in cui è vincolato il mezzo attraverso il quale si propagano le onde è uno solo. Provate a soffiare nel cappuccio di una penna biro: si sente un caratteristico fischio, sempre uguale. Indipendentemente dal modo in cui si soffia, la frequenza del suono percepito è sempre la stessa: varia l'intensità del suono, ma non la frequenza. Il suono prodotto può variare solo se varia (apprezzabilmente) la forma del cappuccio (ad esempio, se si usa il cappuccio di un pennarellone). Il motivo sta nel fatto che il suono è prodotto da un'onda longitudinale che si propaga nell'aria: quando si soffia nel cappuccio l'aria sposta quella presente comprimendola in direzione del fondo del cappuccio. Quando l'onda di compressione giunge sul fondo non può più comprimere altra aria, che non c'è, dunque è costretta a tornare indietro per riflessione. Il risultato è che si produce un'onda di compressione verso il fondo del cappuccio che interferisce con l'onda riflessa da questo che si dirige verso la bocca. Fin quando si continua a soffiare, mentre sul fondo del cappuccio l'ampiezza dell'onda risultante dalla somma di queste due deve rimanere nulla (perché non c'è alcuna

compressione), in prossimità della bocca l'onda deve assumere la massima ampiezza, visto che è quello il punto in cui la compressione, prodotta dal soffio, è massima.

Di nuovo, la somma delle due onde (quella che parte dalla bocca e quella che parte dal fondo) si scrive come

$$y(x, t) = -2A \sin(kx) \cos(\omega t), \quad (7.36)$$

e abbiamo che $y(0, t) = 0$ per ogni t (avendo scelto per $x = 0$ la coordinata del fondo del cappuccio), mentre $y(L, t) = -2A \sin(kL) \cos(\omega t)$. Al variare di t quest'ampiezza dev'essere la massima possibile e quindi si deve avere che

$$kL = (2m + 1) \frac{\pi}{2} \quad (7.37)$$

con m intero, cioè che

$$\frac{2\pi}{\lambda} = (2m + 1) \frac{\pi}{2L} \quad (7.38)$$

e anche in questo caso λ non può avere un valore qualunque, ma deve per forza assumere un valore dato dalla relazione

$$\lambda = \frac{4L}{2m + 1}. \quad (7.39)$$

La massima lunghezza d'onda permessa si ha per $m = 0$ e vale $\lambda = 4L$. Questa lunghezza d'onda corrisponde alla frequenza dell'armonica fondamentale. Le armoniche di ordine superiore hanno una lunghezza d'onda più breve.

Effetti del moto sulle onde

PREREQUISITI: Onde

Le onde sono il prodotto di una qualche sollecitazione che si propaga in un mezzo, che si può considerare fermo. La sorgente delle onde e un osservatore possono essere in moto rispetto a questo mezzo e questo fa sì che le onde subiscano fenomeni di alterazione delle loro caratteristiche.

In questo capitolo studiamo quel che accade a un'onda quando la sorgente o l'osservatore (o entrambi) sono in moto rispetto al mezzo nel quale si propagano le onde.

8.1 L'effetto Doppler

Le onde che si propagano in un mezzo danno luogo a un fenomeno abbastanza peculiare noto con il nome di **effetto Doppler**¹. L'effetto consiste nell'alterazione della frequenza di un'onda misurata da un osservatore o emessa da una sorgente in moto.

Immaginiamo una sorgente di onde (per esempio di onde sonore, come un motore) in avvicinamento con velocità u a un osservatore fermo (a lato della strada). Al tempo $t = 0$ la sorgente emette un primo fronte d'onda che inizia a propagarsi verso l'osservatore a velocità v . La lunghezza d'onda λ e la frequenza ν dell'onda sono legate alla velocità di propagazione v dalla relazione

$$v = \lambda \nu. \quad (8.1)$$

Trascorso un periodo T la sorgente emette un secondo fronte d'onda. In quest'istante la sorgente si è spostata di uT dalla posizione iniziale, mentre il

primo fronte d'onda ha viaggiato per una distanza pari a vT . Di conseguenza la distanza tra i due fronti d'onda è

$$\lambda' = (v - u) T. \quad (8.2)$$

Ma la distanza tra due fronti d'onda per definizione è la lunghezza d'onda, che quindi all'osservatore che la misura appare più piccola rispetto a quella che avrebbe quando la sorgente è ferma. La frequenza ν' percepita dall'osservatore si scrive dunque come

$$\nu' = \frac{v}{\lambda'} = \frac{v}{v - u} \nu. \quad (8.3)$$

È facile vedere che, nel caso in cui la sorgente si allontana dall'osservatore, la frequenza percepita da quest'ultimo è invece

$$\nu' = \frac{v}{v + u} \nu. \quad (8.4)$$

Potete facilmente sperimentare questa situazione mettendovi a bordo strada e ascoltando i rumori prodotti dai mezzi che passano: si ode una caratteristica modulazione del suono prodotto che, quando il mezzo si avvicina, diventa sempre più acuto per diventare sempre più grave quando si allontana. Chiaramente l'effetto è tanto più evidente quanto maggiore è la velocità dei mezzi, perciò se siete in autostrada o a una gara di Formula 1 l'effetto Doppler si sente meglio.

Se è l'osservatore a muoversi quando la sorgente è ferma allora accade qualcosa di diverso: se l'osservatore si avvicina con velocità u alla sorgente di onde, vede queste ultime muoversi verso di lui a una velocità pari a $v' = v + u$ (il tempo impiegato a vedere due fronti d'onda successivi diminuisce perché nel tempo impiegato da un fronte a coprire la distanza

¹Dal nome di **Christian Doppler**.

che lo separa dal precedente, che vale T , l'osservatore si è avvicinato al secondo riducendo il tempo necessario per esserne raggiunto). La frequenza percepita dall'osservatore è dunque

$$\nu' = \frac{v'}{\lambda} = \frac{v+u}{\lambda} = \frac{v}{\lambda} + \frac{u}{\lambda} = \nu \left(1 + \frac{u}{v}\right). \quad (8.5)$$

È evidente che quando l'osservatore si allontana dalla sorgente la relazione che lega ν' a ν diventa

$$\nu' = \left(1 - \frac{u}{v}\right) \nu. \quad (8.6)$$

Questo fenomeno è più difficile da sperimentare perché si verifica quando voi siete in moto e la sorgente è ferma: dovrete trovarvi a passare con la vostra auto a gran velocità nei pressi di un luogo in cui c'è un concerto o una sirena che suona.

In altre parole le onde misurate da un osservatore in moto rispetto a una sorgente cambiano frequenza in modo caratteristico, secondo lo stato di moto di sorgente e osservatore. È utile considerare alcuni casi limite. Quando $u = 0$ la frequenza non deve cambiare e in effetti è così. Se $u = v$, nel caso in cui sia l'osservatore a muoversi abbiamo

$$\nu' = \left(1 \pm \frac{u}{v}\right) \nu \quad (8.7)$$

che può valere $\nu' = 0$ o $\nu' = 2\nu$. Il primo caso corrisponde a quello in cui l'osservatore si allontana dalla sorgente alla stessa velocità con cui viaggiano le onde. È evidente che in questo caso le onde *inseguono* l'osservatore senza mai raggiungerlo e dunque la frequenza misurata è nulla. Nel caso opposto, trascorso un tempo t tale da far avvicinare l'osservatore di metà lunghezza d'onda verso la sorgente, nello stesso tempo il fronte d'onda si è mosso di altrettanto verso l'osservatore e il risultato è che il periodo dell'onda si dimezza (e quindi la frequenza raddoppia).

Quando invece è la sorgente a muoversi, in un caso otteniamo $\nu' = \nu/2$, mentre nell'altro si ottiene $\nu' \rightarrow \infty$. Il primo caso corrisponde al caso in cui la sorgente si allontana dall'osservatore: in questo caso i fronti d'onda risultano più lontani e il periodo si

allunga (raddoppia), quindi la frequenza si riduce. Quando la sorgente si muove verso l'osservatore alla velocità dell'onda ogni volta che emette un fronte d'onda lo fa avendo raggiunto quelli emessi in precedenza, perciò la frequenza è infinita nel senso che la distanza tra un fronte d'onda e il successivo diventa nulla.

L'effetto Doppler per la salute

L'effetto Doppler si usa in numerose applicazioni tecnologiche. Una di queste è l'**ecografia Doppler** o **ecodoppler**. Si tratta di un esame diagnostico che il medico esegue per misurare la velocità con la quale il sangue scorre nei vasi. Il medico ecografista applica al paziente un altoparlante (la sonda ecografica) che produce suoni non udibili dall'orecchio umano. Sulla sonda è montato anche un microfono che rileva le onde riflesse dai tessuti (si ha una riflessione ogni volta che l'onda incontra la superficie di separazione tra due tessuti di diversa densità). Anche il sangue riflette le onde, ma a differenza dei tessuti non è fermo: se il sangue si sta allontanando dalla sonda *vede* una frequenza minore rispetto a quella prodotta dall'altoparlante. Le onde riflesse poi sono prodotte da una sorgente in allontanamento e quindi per la sonda hanno una frequenza ancora minore. Misurando la differenza di frequenza rispetto alle onde inviate si può conoscere la velocità (e il verso) del sangue.

Anche le onde radio sono soggette a effetto Doppler (**relativistico** però) e la polizia le usa per misurare la velocità dei veicoli con i cosiddetti **autovelox**. Onde radio di frequenza ν prodotte da una sorgente ferma sono inviate sulle auto in moto lungo una strada: quando giungono sull'auto questa *vede* le onde con una frequenza diversa ν' e le riflette diventando così una sorgente in moto di onde di frequenza ν' . Al radar giungono così onde radio emesse alla frequenza ν' ulteriormente alterate dal fatto che la sorgente (l'auto) è in moto. Misurando la frequenza ν'' dell'onda così rilevata dai sensori del radar si risale alla velocità dell'auto.

Quando sia la sorgente che l'osservatore sono in moto rispetto al mezzo, l'effetto Doppler complessivo è quello che si ottiene combinando i due effetti perciò

$$\nu' = \left(1 \pm \frac{v}{u}\right) \frac{v}{v \pm u} \nu = \frac{u \pm v}{v \pm u} \nu. \quad (8.8)$$

È utile a questo punto porsi la seguente domanda: se l'osservatore si muove con velocità u verso la sorgente, gli effetti che ci si aspetta di vedere dovrebbero essere gli stessi di quando è la sorgente a muoversi con velocità u verso l'osservatore! In fondo, il moto è sempre relativo! Eppure, nel caso dell'effetto Doppler sembra che non sia così: le relazioni che danno la frequenza percepita dell'onda nei due casi sono diverse! Come mai? La risposta a questa domanda è semplice: l'osservazione è semplicemente sbagliata. Quel che è in moto relativo nell'effetto Doppler non è la coppia sorgente-osservatore, ma la coppia onda-osservatore. L'onda si muove, rispetto al mezzo che è fermo, a velocità v sia che la sorgente sia ferma che nel caso in cui la sorgente sia in moto. Quando l'osservatore corre verso la sorgente, la velocità apparente delle onde è diversa, ma non è così quando è la sorgente a muoversi (vedi il Video (8.1)).

8.2 Effetto Doppler Relativistico

La luce è un'onda elettromagnetica, come le onde radio, i raggi infrarossi e quelli ultravioletti, i raggi X, i raggi γ (vedi 9.2). Anch'essa dunque è soggetta all'effetto Doppler, ma con un'importante differenza. Gli esperimenti provano che la velocità della luce è sempre la stessa, qualunque sia lo stato di moto dell'osservatore. Per quanto possa sembrare strano, se uno si muove molto rapidamente (ad esempio, a $u = 200\,000$ km/s) andando incontro a un raggio di luce, che viaggia a $c = 300\,000$ km/s, non vede quest'ultimo venirgli incontro a velocità $U = u + c = 500\,000$ km/s, ma sempre a velocità c . Analogamente, se s'insegue un raggio di luce a velocità molto alta u , la velocità con la quale lo si

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 8.1 Nel film di **Sergio Corbucci** "Lo smemorato di Collegno" (1962), Totò è un disgraziato che ha perso la memoria ed è stato rinchiuso in manicomio. In questa scena **Erminio Macario** prova a convincere Totò che agitare la testa davanti a un ventaglio fermo è la stessa cosa che agitare il ventaglio davanti alla testa ferma, perché il moto relativo è lo stesso. Purtroppo quello che conta non è il moto relativo tra testa e ventaglio, ma quello tra aria (mossa dal ventaglio) e testa, proprio come nell'effetto Doppler nel quale ciò che conta non è il moto relativo tra sorgente e osservatore, ma quello tra osservatore e fronti d'onda [<https://www.youtube.com/watch?v=f-j5zPP158o&feature=youtu.be>].

vede allontanarsi non è $U = c - u$, ma sempre c ! Quest'osservazione sperimentale portò alla formulazione della **relatività ristretta**.

Quando la sorgente luminosa si avvicina o si allontana da un osservatore fermo, la frequenza percepita da quest'ultimo è data dalla relazione che abbiamo trovato per tutte le onde:

$$\nu' = \frac{c}{c \pm u} \nu. \quad (8.9)$$

Nessuna sorgente luminosa può muoversi a velocità maggiori o uguali a c perciò $c - u$ è sempre un numero positivo e $c + u$ è sempre minore di $2c$. Se è l'osservatore a muoversi verso la sorgente, nel caso della luce questi non vede più l'onda propagarsi a una velocità $v' = c + u$; la velocità di propagazione delle onde elettromagnetiche è sempre c , ma la lunghezza d'onda di queste diminuisce perché l'os-

servatore incontra due fronti d'onda successivi in tempi piú brevi per cui

$$\lambda' = \lambda - uT \quad (8.10)$$

per cui

$$\nu' = \frac{c}{\lambda'} = \frac{c}{\lambda - uT}. \quad (8.11)$$

Essendo $\lambda = \frac{c}{\nu}$ e $T = 1/\nu$, sostituendo e manipolando un po' l'equazione si ottiene

$$\nu' = \frac{c}{\frac{c}{\nu} - \frac{u}{\nu}} = \frac{c}{c - u} \nu. \quad (8.12)$$

È evidente che, nel caso in cui l'osservatore si allontana, si avrà il segno meno a denominatore e in definitiva

$$\nu' = \frac{c}{c \pm u} \nu. \quad (8.13)$$

secondo il verso del moto relativo. La relazione che c'è tra le frequenze emesse e quelle misurate è la stessa nei due casi perché nel caso della luce il moto è effettivamente relativo: non è possibile stabilire chi si muove e chi sta fermo!

L'effetto Doppler relativistico altera la frequenza della luce emessa da oggetti in rapido movimento rispetto a noi. Dal momento che noi percepiamo le onde luminose di lunghezza d'onda attorno ai 700 nm come luce rossa e quelle con $\lambda \simeq 400$ nm come blu, una sorgente di luce di un determinato colore che si allontana da noi sarebbe percepita come una sorgente di colore diverso, con lunghezza d'onda maggiore. Da qui il nome di **spostamento verso il rosso** o **red shift** che si dà al fenomeno che si verifica osservando galassie molto lontane dalla nostra (vedi Par. 8.4).

8.3 L'effetto Čerenkov

Se la velocità con la quale si muove la sorgente è maggiore di quella con la quale l'onda si propaga, allora il fronte d'onda successivo al primo è emesso in un punto tale per cui nel tempo $t = T$ l'onda si

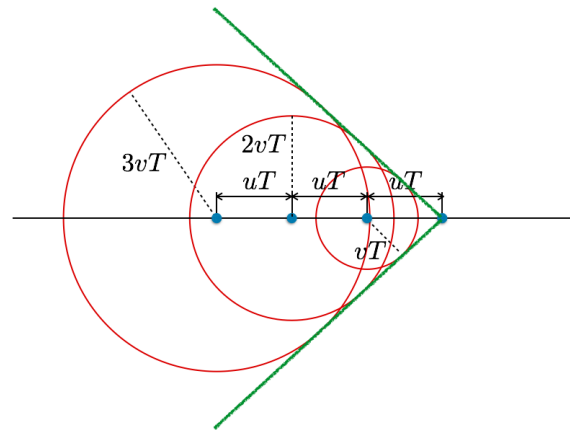


Figura 8.2 Una sorgente di onde sferiche si muove a velocità costante u verso destra emettendo onde che si muovono con velocità v . Il fronte d'onda che ne risulta è quello rappresentato in verde.

è propagata per una distanza inferiore a quella percorsa dalla sorgente. Il secondo fronte d'onda quindi è emesso quando la sorgente è al di là del primo fronte d'onda. Quando la sorgente emette il terzo fronte d'onda, il primo si è propagato per una distanza pari a $d_1 = 2vT$, il secondo per una distanza $d_2 = vT$, mentre la sorgente si trova a una distanza $d_S = 2uT > 2vT$. Nel caso di onde bidimensionali (come quelle sull'acqua) Possiamo dunque rappresentare il primo fronte d'onda come una circonferenza di raggio $2v$ centrata nell'origine, perché emessa da tale punto, e il secondo come una circonferenza di raggio v centrata in un punto che dista uT dall'origine. Dopo un ulteriore tempo T il primo fronte d'onda si è propagato per una distanza di $3vt$, il secondo di $2vT$ a partire da un punto che dista uT dal centro del primo fronte d'onda e il terzo per una distanza vT a partire da un punto che dista uT dalla sorgente del secondo fronte, come nella Fig. 8.2.

L'involuppo di tutti i fronti d'onda che si muovono tutti insieme in una direzione non è altro che un'altra onda costituita dalla somma delle singole onde. Nella Figura 8.2 è rappresentato con una riga verde e ha la forma di un cono (o di una sua sezio-

ne se l'onda è bidimensionale). L'apertura del cono, cioè l'angolo con il quale le onde piane che costituiscono l'involuppo appaiono provenire, dipende dalla velocità della sorgente: infatti, quando la sorgente ha percorso un tratto lungo uT , il fronte d'onda è avanzato di un tratto vT . Se chiamiamo α l'angolo formato tra la direzione di volo della sorgente e il fronte d'onda piano provocato da questo fenomeno abbiamo che

$$uT \sin \alpha = vT \quad (8.14)$$

e quindi

$$\sin \alpha = \frac{v}{u}. \quad (8.15)$$

Più è alta la velocità della sorgente u , più è piccolo il seno dell'angolo (e quindi l'angolo). Questo è il fenomeno che provoca l'apparire di una scia dietro le imbarcazioni che si muovono più rapidamente di quanto facciano le onde sull'acqua (Figura 8.3) [6].

Nel caso della luce, la velocità di questa in un mezzo d'indice di rifrazione n è $v = c/n$ con $c = 3 \times 10^8 \text{ m s}^{-1}$ la velocità della luce nel vuoto. Si ottiene quindi che, se una sorgente luminosa viaggia a una velocità $u > c/n$ in quel mezzo, il fronte d'onda della luce appare conico. In questo caso la luce emessa forma un fronte di semiapertura θ pari a

$$\sin \theta = \frac{c}{nu}. \quad (8.16)$$

La luce appare provenire da una direzione che è perpendicolare a quella lungo la quale si distribuisce il fronte d'onda (che è la direzione del raggio che abbiamo preso in considerazione per calcolare θ) e perciò forma un angolo rispetto alla direzione di volo della sorgente pari a

$$\theta' = \pi - \frac{\pi}{2} - \theta = \frac{\pi}{2} - \theta \quad (8.17)$$

(abbiamo usato la proprietà dei triangoli secondo la quale la somma degli angoli interni vale sempre π). Il coseno di quest'angolo è uguale al seno di θ e perciò

$$\cos \theta' = \frac{c}{nu} \quad (8.18)$$



Figura 8.3 La scia lasciata da quest'anatra che nuota sulla superficie di un lago è provocata dall'effetto discusso nel paragrafo: l'animale si sposta più velocemente di quanto non faccia l'onda che si propaga sull'acqua. Questa è provocata dallo spostamento del liquido prodotto dall'anatra che, spostandosi, spinge verso il basso una porzione di acqua, la quale, essendo incompressibile, spinge verso l'alto quella circostante. Fino a valori di velocità u della sorgente inferiori a $gL/2\pi$, con g pari all'accelerazione di gravità e L alla lunghezza dello scafo, il fenomeno si descrive secondo un modello dovuto a Lord Kelvin. Al di sopra di questa soglia l'apertura del cono dipende da u .

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 8.4 Nell'esperimento illustrato in questo filmato si dimostra come l'apertura del cono formato dalla scia di un oggetto che si muove sull'acqua dipende dalla sua velocità. Per gentile concessione dei Proff. M. Rabaud e F. Moisy [<http://www.fast.unpsud.fr/~moisy/wake/>] [6].

Piú veloce della luce

Anche la luce è un'onda di natura elettromagnetica. Le onde elettromagnetiche sono emesse dalle cariche elettriche in moto. Una particella elettricamente carica può muoversi, in un mezzo come l'acqua, con una velocità superiore a quella della luce nello stesso mezzo (quella che non si può superare è la velocità della luce nel vuoto!). In questo caso le onde elettromagnetiche emesse dalle cariche accelerate si sommano in un fronte d'onda conico che produce un segnale elettromagnetico che si propaga come un'onda sferica di lunghezza d'onda tale, in certi casi, da corrispondere a quella della luce blu. La luminescenza azzurrina (vedi Fig. 8.6) che si vede attorno alle barre di combustibile di una centrale nucleare quando sono immerse in acqua è proprio prodotta dall'effetto Čerenkov provocato dagli elettroni emessi dall'uranio che nell'acqua si muovono piú rapidamente della luce. Lo stesso fenomeno si sfrutta per rivelare il passaggio di particelle cariche all'interno di serbatoi d'acqua (Fig. 8.5) sulle cui pareti sono montati dei fotomoltiplicatori (sensori sensibili alla luce): il cono di luce prodotto dalle particelle che attraversano l'acqua, interseca la parete del contenitore e *accende* i fotomoltiplicatori disposti lungo un'ellisse.

è la direzione di propagazione della luce relativa alla direzione di volo della sorgente. Il fenomeno si sfrutta per rivelare il passaggio di particelle elettricamente cariche negli esperimenti. Queste particelle, infatti, passando attraverso un materiale ne *polarizzano* le molecole, le quali, depolarizzandosi, emettono *radiazione elettromagnetica* che, in certi casi, può risultare visibile. Il fenomeno prende il nome di **effetto Čerenkov**. È come se ogni punto della traiettoria percorsa dalla particella si comportasse come una sorgente di *onde luminose* che si muovono, nel mezzo, a una velocità inferiore a quella con la quale la sorgente si sposta. Il risultato è che la luce prodotta dal fenomeno assume la forma di un cono. Se le particelle che attraversano il mez-

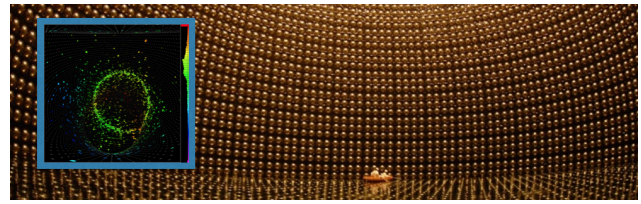


Figura 8.5 L'esperimento **Super-Kamiokande**, in Giappone, consiste in una enorme vasca riempita d'acqua le cui pareti sono ricoperte di tubi fotomoltiplicatori. Una particella carica abbastanza veloce, attraversando l'acqua, produce un cono di luce Čerenkov rivelata dai fotomoltiplicatori. L'intersezione tra il cono Čerenkov e la parete del rivelatore è (quasi) una conica, che il piú delle volte è un'ellisse. Nell'insero si vede l'immagine della luce Čerenkov ricostruita dai computer dell'esperimento.

zo sono tante l'intensità della luce può addirittura essere visibile a occhio nudo.

8.4 Il red shift delle galassie

Come si dimostra nel Capitolo 9 la luce è un'onda, il cui colore dipende dalla lunghezza d'onda λ . Se s'invia luce bianca su un prisma di vetro, dalla faccia opposta di questo emerge un fascio di luce allargato rispetto a quello incidente, che mostra tutti i colori dell'arcobaleno. Il fenomeno s'interpreta facilmente anche alla luce della teoria ondulatoria della luce (Paragrafo 6.2): è sufficiente assumere che la luce bianca sia una **sovrapposizione** di onde di frequenza e lunghezza d'onda diverse, la cui velocità di propagazione nel vetro dipende dalla lunghezza d'onda. In questo modo le onde che prima d'incontrare il prisma viaggiavano parallele l'una all'altra, si **disperdono** attraversando il vetro perché ciascuna è rifratta a un angolo diverso secondo la propria lunghezza d'onda. Quando emerge dal vetro, dunque, la luce rossa risulta deviata in maniera diversa

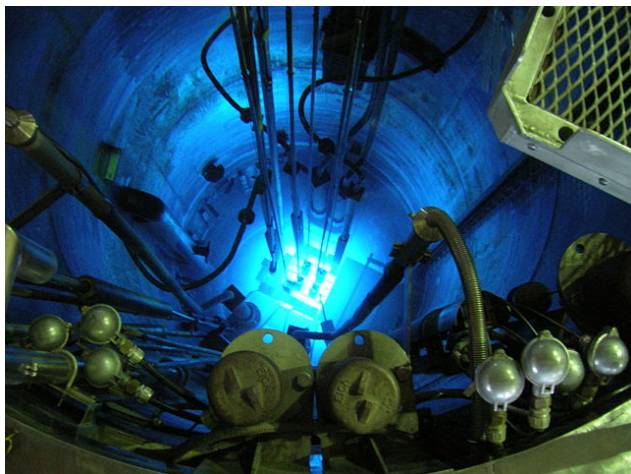


Figura 8.6 La luminescenza blu che si vede nella vasca di un reattore nucleare è prodotta dalla luce Čerenkov emessa dagli elettroni prodotti dall'Uranio radioattivo del combustibile, i quali si muovono più rapidamente della luce in acqua (immagine di Pieck Darío).

da quella blu. L'insieme dei colori che si osservano si chiama **spettro** della luce bianca.

Se, prima di entrare nel prisma, la luce bianca passa attraverso un recipiente trasparente nel quale è contenuto un gas o un liquido in cui è presente una particolare sostanza, dallo spettro che si può osservare a valle del prisma *mancano* alcuni colori. Si dice che si osservano le **righe di assorbimento** (il termine **riga** deriva dal fatto che di norma la sorgente di luce bianca ha la forma di un rettangolo lungo e stretto che all'occhio appare come una riga e l'immagine di questa sorgente che si forma dopo l'attraversamento del prisma ha la stessa forma) e s'interpreta il fatto assumendo che la sostanza presente nel recipiente ha la capacità di **assorbire** la luce incidente solo quando la sua lunghezza d'onda assume certi particolari valori. Se le onde con certe frequenze sono assorbite dalla sostanza prima di entrare nel prisma, una volta dispersa la luce i corrispondenti colori mancheranno dallo spettro.

La distribuzione delle righe di assorbimento è caratteristica di ogni specie chimica. Ad esempio, l'idrogeno fa *sparire* le onde di lunghezza d'onda pari

a 434, 486 e 656 nm, corrispondenti a certe tonalità dei colori violetto, blu e rosso. Osservando uno spettro di assorbimento dunque si può ricavare la composizione chimica delle sostanze che la luce ha attraversato prima di essere dispersa. È così, ad esempio, che possiamo stabilire di quali elementi chimici sono composte le stelle: osservandone la luce fatta passare attraverso un prisma e individuando le righe mancanti. La luce che proviene dall'interno della stella infatti ne ha dovuto attraversare l'atmosfera e quindi parte di essa sarà stata assorbita.

Se si osserva lo spettro della luce di galassie lontane si vede qualcosa di strano: le righe a 434, 486 e 656 non ci sono, ma ce ne sono altre che non si osservano mai in laboratorio! Una prima interpretazione di questo fatto si potrebbe dare dicendo che le galassie lontane evidentemente devono essere composte di elementi diversi da quelli presenti sulla Terra. Certo, questo è strano, e se si guarda meglio si scopre che le righe presenti, pur non trovandosi dove ci si aspetta, hanno la stessa distribuzione di quelle che si osservano sulla Terra. Per esempio, se la prima riga è a 480 nm, la seconda è a 537 e la terza a 726. I rapporti tra le lunghezze d'onda delle righe della galassia e quelle delle righe dell'idrogeno a Terra sono tutti uguali a circa 1.11: è come se le lunghezze d'onda delle righe dell'idrogeno fossero state tutte *stirate* di un fattore 1.11, per cui le righe di assorbimento si sono spostate tutte verso le lunghezze d'onda maggiori (cioè verso il rosso, da cui il nome **spostamento verso il rosso** o **red shift** per il fenomeno).

Una possibile spiegazione allora consiste nel pensare che lo spostamento delle righe sia dovuto all'**effetto Doppler** per il quale l'equazione (8.12) prevede che

$$\nu' = \frac{c}{\frac{c}{\nu} - \frac{u}{\nu}} = \frac{c}{c - u} \nu, \quad (8.19)$$

quando sorgente e osservatore si avvicinano l'uno all'altro a velocità u . Riscrivendo la formula per le lunghezze d'onda si trova

$$\nu' = \frac{c}{\lambda'} = \frac{c}{c - u} \frac{c}{\lambda} \quad (8.20)$$

da cui

$$\lambda' = \frac{c - u}{c} \lambda. \quad (8.21)$$

Poiché nel nostro caso $z = \frac{\lambda'}{\lambda} \simeq 1.1$, dev'essere

$$\frac{c - u}{c} = 1 - \frac{u}{c} = z. \quad (8.22)$$

Visto che $z > 1$, u dev'essere negativa e questo significa semplicemente che la sorgente (la galassia) si sta allontanando da noi (o che noi ci stiamo allontanando dalla sorgente, ma questo non fa differenza) e che la velocità di allontanamento è, in modulo, $u \simeq 0.11c = 0.33 \times 10^8$ m/s.

La diffrazione

PREREQUISITI: Interferenza

Una delle dispute più longeve della fisica ha riguardato la natura della luce. **Newton** pensava [7] che la luce fosse costituita di minuscole particelle, mentre l'olandese **Christiaan Huygens** sosteneva che si dovesse trattare come un'onda¹. Entrambi avevano evidentemente buoni motivi per sostenere l'una o l'altra ipotesi: definire la natura di qualcosa è possibile solamente grazie ai dati sperimentali in nostro possesso. A quanto ne sappiamo la luce può essere riflessa da uno specchio o rifratta da una lente, ma questi due fenomeni si possono verificare tanto con le onde che con le particelle.

Le onde del mare o i suoni, ad esempio, sono riflessi quando incontrano un ostacolo di dimensioni sufficienti. L'eco, in fondo, non è che la riflessione di un suono prodotta da una parete rocciosa: perché si verifichi è necessario che la parete sia a distanza sufficiente a produrre un ritardo apprezzabile, sia abbastanza grande da intercettare una frazione significativa dell'onda sonora e abbastanza liscia da evitare la dispersione dell'onda riflessa in molte direzioni. Anche il materiale di cui è fatta la parete deve possedere proprietà elastiche adeguate a produrre il fenomeno. Allo stesso modo, un'onda luminosa può essere o meno riflessa da una superficie, secondo che questa abbia o meno certe caratteristiche: deve essere liscia, ad esempio, ma questo non basta. Un foglio di carta può essere liscio, ma non

riflette la luce. Deve apparire *lucido* come i metalli². La distanza e le dimensioni trasversali solitamente non hanno alcun effetto sulla possibilità di riflettere la luce, suggerendo che la lunghezza caratteristica (la sua lunghezza d'onda) di un'eventuale onda luminosa sia molto più piccola di quella degli specchi che usiamo.

Le onde del mare cambiano direzione (rifrazione) quando cambia la loro velocità di propagazione: quest'ultima, nel caso delle onde del mare, dipende dalla profondità del fondale. Fino a quando la profondità del fondale è molto maggiore dell'ampiezza dell'onda, questa si propaga in linea retta con una certa velocità. Di conseguenza il fronte d'onda appare piano in alto mare. Quando l'onda si avvicina alla costa, però, il fronte comincia a rallentare nei punti in cui il fondale è più basso. Di conseguenza il fronte si *piega* e l'onda, pur mantenendo un fronte d'onda approssimativamente piano, assume una direzione diversa da quella originale. Come si può facilmente osservare stando in spiaggia, le onde del mare viaggiano sempre grosso modo in direzione della costa, mai parallelamente a questa. Lo stesso fa la luce: si propaga di norma in linea retta, ma se incontra un mezzo nel quale la sua velocità è diversa, la direzione di propagazione cambia in funzione delle caratteristiche del mezzo. In particolare, se la velocità si riduce, la direzione di propagazione tende ad avvicinarsi alla normale alla superficie di separazione tra i due mezzi.

Rifrazione e riflessione, insomma, sono due fenomeni che si possono attribuire ai fenomeni ondulatori, perciò se ne potrebbe concludere che la luce dev'essere rappresentabile come un'onda che si pro-

¹In realtà non ci fu mai una vera disputa in tal senso, perché la fama di Newton all'epoca era infinitamente superiore a quella di Huygens, i cui lavori erano probabilmente ignorati dall'inglese.

²C'è una ragione per questo, che potrete scoprire [qui](#)

paga nella direzione dei *raggi* luminosi. Resterebbe da stabilire la natura dell'onda (cos'è che oscilla nel caso della luce?), ma questo è un problema diverso.

Ma a pensarci bene, rifrazione e riflessione sono fenomeni che si potrebbero verificare anche se la luce fosse composta di particelle! È esperienza piuttosto comune, ad esempio, che se si scaglia una palla su un muro, questa rimbalza, e che l'angolo formato tra la velocità della palla dopo l'urto e la normale alla superficie del muro è uguale a quello formato tra questa direzione e la velocità della palla prima dell'urto: è una banale conseguenza della [conservazione della quantità di moto](#). La riflessione della luce, dunque, si potrebbe spiegare in questo modo.

In modo del tutto analogo, se si assume che la luce sia composta di corpuscoli che si muovono parallelamente l'uno all'altro, quando il fronte avanza formando un certo angolo con la normale alla superficie di separazione tra due mezzi nei quali la velocità di propagazione è diversa, i corpuscoli che giungono prima cambiano velocità prima degli altri (Vedi il Filmato 5.10). Questo provoca il *piegamento* del fronte e l'alterazione della direzione di propagazione. Questo principio è utilizzato, ad esempio, dai mezzi cingolati per sterzare: non potendo ruotare l'asse delle ruote, questi mezzi fanno muovere i cingoli a velocità diverse. Così facendo, uno dei lati del mezzo resta indietro rispetto all'altro e di fatto il mezzo sterza. La rifrazione della luce dunque si può spiegare assumendo che i corpuscoli che la formano, giunti sulla superficie di separazione tra due mezzi, cambiano direzione per il semplice fatto che quelli che arrivano in anticipo rallentano (o accelerano) prima degli altri.

C'è però un fenomeno che si verifica solo con le onde: la **diffrazione**, che consiste nella possibilità che hanno le onde di aggirare gli ostacoli grazie alla capacità di modificare la forma del fronte d'onda e la sua direzione di propagazione.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 9.1 Il suono è provocato da un'onda di compressione delle molecole d'aria che si propaga. In questo video l'animazione delle particelle d'aria che si muovono in seguito alla sollecitazione prodotta da un altoparlante è seguito da un filmato fatto con una tecnica chiamata *Schlieren flow*. Per comprendere bene la propagazione del suono scegliete una particella qualunque e seguitene il moto con l'occhio: vedrete che ciascuna particella si muove solo attorno a una posizione di equilibrio e che quel che si propaga è la sollecitazione prodotta dalla collisione di ciascuna particella con una sua vicina. [Il video integrale, che illustra anche la tecnica di ripresa, è all'indirizzo <https://www.youtube.com/watch?v=px3oVGXr4mo>]. Per gentile concessione del Prof. Michael Hargather, con la collaborazione del Dr. Gary Settles della Penn State University.

9.1 Sperimentiamo la diffrazione

Alcuni fenomeni di diffrazione si possono sperimentare facilmente: ad esempio, se ci si pone dietro un'ostacolo di dimensioni relativamente piccole frapposto tra noi e una sorgente sonora (ad esempio, una colonna), il suono si percepisce lo stesso perché le onde sonore riescono ad *aggirare* l'ostacolo. È facile capire come avviene, a livello microscopico, il fenomeno: le onde di pressione che costituiscono il suono sono costituite di fluttuazioni di densità di particelle che si muovono in tutte le direzioni (vedi

filmato 9.1). In presenza di un ostacolo, le particelle che lo urtano, oltre alla forza di compressione del suono, subiscono la forza di reazione dell'ostacolo che ne modifica la direzione del moto, alterando localmente la forza di compressione che produce il suono. Il risultato di molte interazioni di questo tipo è che il suono (l'onda, non le particelle) aggira l'ostacolo e noi possiamo ascoltarlo anche stando dietro di questo. È del tutto evidente che se il suono fosse formato di particelle che si muovono nella direzione di propagazione, come nel caso del vento, frapponendo un ostacolo tra la sorgente e l'ascoltatore, questo moto sarebbe impedito e dunque non si potrebbe ascoltare nulla al di là dell'ostacolo.

Allo stesso modo, se l'onda passa attraverso un'apertura praticata su un ostacolo di grandi dimensioni, al di là di questo l'onda si propaga in maniera diversa rispetto a quanto faceva prima d'incontrare l'ostacolo. Se, ad esempio, vi chiudete in una stanza, non riuscite a percepire i rumori esterni (a meno che non siano particolarmente intensi, perché i suoni, pur fortemente attenuati, si trasmettono anche attraverso i muri, le cui particelle vibrano come quelle d'aria, ma con minore ampiezza). Ma se si apre una porta o una finestra, anche se non ci si mette esattamente in corrispondenza dell'apertura, si percepisce un suono intenso, anche se ci si mette dietro un angolo. Il motivo è che l'onda sonora, giungendo in prossimità della porta o della finestra, comincia a comprimere le molecole di aria presenti in corrispondenza del varco che a loro volta cominciano a comprimere le molecole adiacenti, generando un moto complessivo al di là dell'ostacolo. Ciascuna molecola di aria si comporta come una sorgente di onde (è quello che si chiama **principio di Huygens**, dal nome di [Christiaan Huygens](#) che lo formulò per primo).

Il fronte d'onda sonoro che giunge in corrispondenza di una porta provoca l'emissione di un fronte al di là di questa, che si propaga con onde semisferiche all'interno del locale e giunge al nostro orecchio anche se non ci troviamo in corrispondenza della porta da cui entra il suono. Il suono si può propagare in tutta la stanza anche se l'apertura praticata sul muro è sottilissima, grazie allo stesso fenomeno.

Le particelle di aria presenti lungo lo stretto varco si comportano ciascuna come sorgenti di onde sferiche che quindi si propagano in tutta la stanza. Per isolare acusticamente una stanza dunque è necessario sigillarla ermeticamente.

Un altro esempio piuttosto evidente di come agisca la diffrazione è quello mostrato in Figura 9.2, in cui si vede una foto satellitare (presa da [Google Maps](#) nel mese di luglio 2014) della costa vicino a Fiumicino, in provincia di Roma. Davanti alla spiaggia sono state sistemate alcune *dighe* per arginare il moto ondoso e limitare l'erosione della spiaggia da parte delle onde del mare. Tra una diga e l'altra c'è un piccolo varco. Infrangendosi sulle dighe, le onde marine sono riflesse, ma attraverso i varchi possono continuare a propagarsi. Nel farlo, però, la forma del fronte d'onda, inizialmente piana, si modifica e diventa approssimativamente circolare. Di fatto il varco agisce come una sorgente di onde circolari (avendo dimensioni relativamente piccole) che si propagano da entrambi i lati. Dal lato del mare l'onda si somma con quella incidente del mare e l'effetto è trascurabile; dall'altro lato invece l'onda si propaga quasi liberamente ed erode la spiaggia dandole la caratteristica forma che coincide con quella del fronte che si allarga.

Anche in questo caso, se le onde del mare fossero rappresentabili come particelle in moto verso la costa, in corrispondenza dei varchi vedremmo la costa erosa perpendicolarmente alla diga e non a semicerchio.

9.2 Definiamo la natura della luce

Per capire la natura della luce, dunque, è necessario eseguire un esperimento col quale si possa stabilire se la luce dà origine a fenomeni di diffrazione o meno. Un modo per verificare se la luce è un'onda potrebbe consistere nel produrre una **figura di diffrazione** simile alla Figura 9.2: basterebbe inviare un fascio di luce sul muro interponendo uno schermo



Figura 9.2 La costa italiana in prossimità del comune di Fiumicino (Roma) ha la forma dei fronti d'onda marini prodotti dalla diffrazione provocata dalle aperture tra una diga e l'altra, poste a salvaguardia della costa nel mare.

sul quale sono praticate alcune aperture³. Se al di là dello schermo osserviamo le immagini delle aperture possiamo concludere che la luce è formata di particelle: queste, propagandosi perpendicolarmente allo schermo, si dovrebbero fermare quando ne sono ostacolate, mentre dovrebbero proseguire in linea retta in corrispondenza delle aperture. Se, per esempio, le aperture sono formate da fori circolari, ci aspetteremmo di vedere dei cerchi luminosi formarsi sul muro. Nel caso in cui l'apertura sia un taglio verticale ci aspetteremmo di vedere una **riga** luminosa sul muro. Se, al contrario, osservassimo luce in punti del muro non corrispondenti alle aperture praticate sullo schermo, dovremmo concluderne che la luce ha una natura ondulatoria.

A dire il vero la conclusione non potrà essere netta, nel caso in cui non si osservi alcun fenomeno diffrattivo. Pensate al caso delle dighe di fronte a Fiumicino: aumentando progressivamente l'ampiezza delle aperture, a un certo punto succede che le dighe non provocano più alcun effetto sulle onde che si frangono sulla costa, che continueranno ad andare dritte! Al limite le dighe potrebbero essere larghe

solo pochi centimetri e distare tra loro molti metri. Solo se le aperture delle dighe sono abbastanza piccole si osserva il fenomeno. Come al solito i termini **piccolo** e **grande** in fisica assumono significato solo in relazione a qualche grandezza fisica omogenea. Nel caso specifico, in cui la diffrazione dipende dall'ampiezza dell'apertura praticata sulla diga, che dimensionalmente è una **lunghezza**, è chiaro che per il verificarsi della diffrazione è essenziale che quest'ampiezza deve essere piccola rispetto a un'altra lunghezza che deve caratterizzare l'onda, che non può che essere la sua **lunghezza d'onda**. In definitiva la diffrazione si verifica solo se l'ampiezza dell'apertura praticata sulla diga è abbastanza piccola rispetto alla lunghezza d'onda delle onde del mare.

Se quindi non vedessimo diffrazione da onde luminose potrebbe semplicemente essere che la lunghezza d'onda della luce sia troppo piccola rispetto alle dimensioni delle fenditure. Potremmo solo mettere un limite superiore alla eventuale lunghezza d'onda.

Abbiamo già osservato che non sembrano esserci limiti alla possibilità per uno specchio di riflettere efficacemente la luce, al contrario di quel che accade nel caso delle onde sonore, che si riflettono solo se l'ostacolo ha grandi dimensioni. Questo significa sicuramente che la lunghezza d'onda deve essere molto piccola: più piccola delle dimensioni con le quali siamo in grado di fabbricare gli specchi. Se dunque vogliamo tentare un esperimento di diffrazione dobbiamo usare uno schermo con fori di dimensioni submillimetriche.

Quello che ci possiamo aspettare è che se la luce fosse composta di corpuscoli, inviando un fascio luminoso su uno stretto taglio praticato su uno schermo, al di là dello schermo dovremmo aspettarci un'immagine abbastanza fedele del taglio. Possiamo aspettarci qualche *sbavatura* ai bordi perché magari i corpuscoli si muovono con velocità tali per cui quelli che passano vicino al bordo, urtando su quest'ultimo, possono essere deviati di un angolo più o meno grande, ma di sicuro ci aspettiamo il comparire di una sola immagine della fenditura. Se invece mettiamo un ostacolo davanti al fascio ci aspettiamo di vedere l'*ombra* dell'ostacolo qualunque siano

³Quest'esperimento fu proposto per la prima volta da Augustin Fresnel [8]

le dimensioni di questo e la distanza tra l'ostacolo e la superficie sulla quale si forma l'ombra.

Se, al contrario, la luce è un'onda, dobbiamo aspettarci il comparire di effetti diffrattivi quando le dimensioni degli ostacoli o delle fenditure sono comparabili con quelle della lunghezza d'onda della luce. In questo caso, facendo passare la luce attraverso una stretta fenditura ci dobbiamo aspettare il formarsi di un'immagine simile a quella della Figura 9.2. Inoltre, in presenza di un ostacolo abbastanza piccolo, ci aspettiamo di non vederne l'ombra.

9.3 La matematica della diffrazione

Per comprendere meglio il fenomeno della diffrazione conviene fare qualche conto per capire cosa ci si deve attendere nei diversi casi esaminati. Dal momento che la diffrazione si verifica perché ciascun punto del mezzo nel quale si propaga un'onda diventa a sua volta sorgente di onde, in certi casi la diffrazione è accompagnata dal fenomeno dell'interferenza, provocata dal sommarsi delle molte onde prodotte. Se non ve la sentite di fare i conti usando la trigonometria, potete comunque rendervi conto di cosa accade, anche se in maniera qualitativa, usando un'applicazione **GEOGEBRA**⁴.

9.3.1 Diffrazione da fenditura

Immaginiamo di aver praticato su uno schermo una fenditura molto stretta: molto più stretta della lunghezza d'onda dell'onda incidente. In questo caso possiamo trattare la fenditura come una sorgente di onde puntiforme. Il fronte d'onda al di là dello schermo avrà una forma sferica se l'onda si propaga nelle tre dimensioni (come il suono) o circolare se la propagazione avviene su una superficie (come le onde su un liquido). Se da un lato dello schermo arrivano onde piane, dall'altro si osserva un'onda sferica o circolare propagarsi partendo da un pun-

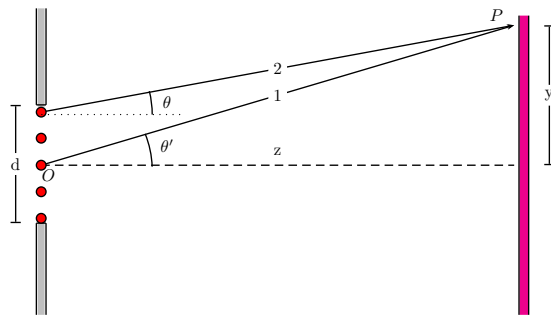


Figura 9.3 L'onda che giunge in un punto distante z dalla fenditura nella direzione di propagazione dell'onda incidente e y dal punto medio della fenditura in direzione ortogonale è il risultato della sovrapposizione di tutte le onde generate da ciascuno dei punti della fenditura. Nella figura ne sono indicati cinque.

to che corrisponde alla posizione della fenditura (è quel che si vede, ad esempio, in Fig. 9.2).

Non abbiamo molto da calcolare in questo caso. L'onda appare provenire da ogni direzione uscente dalla fenditura. Tutti gli angoli da $-\pi$ a π rispetto alla normale allo schermo sono permessi.

Se consideriamo una fenditura di ampiezza d non troppo piccola (in maniera da non poter considerare perfettamente puntiforme la sorgente), accade qualcosa di curioso: l'onda si propaga sempre al di là dello schermo, ma sembra assente se si guarda in direzione di certi angoli.

Ponendo lo schermo a distanza z dal punto d'osservazione (Fig. 9.3), da ciascun punto della fenditura (alcuni di questi sono marcati in rosso nella figura) partirà un treno di onde di lunghezza d'onda λ e in fase tra loro. Per fissare le idee facciamo l'ipotesi che la luce sia un'onda, per cui illuminiamo la fenditura con un laser (in modo tale che la luce sia **monocromatica**) e osserviamo il formarsi dell'immagine su un muro. Su uno stesso punto del muro, dunque, confluiranno gli effetti di tutte le onde emesse da ogni punto della fenditura.

⁴All'indirizzo <http://tube.geogebra.org/student/m766777>

In particolare, dal punto piú in alto della fenditura nella Figura 9.3, parte un'onda sferica in fase con quella incidente che dunque si propaga in tutte le direzioni. Consideriamo la direzione che forma un angolo θ con la normale alla fenditura. L'onda giungerà, propagandosi lungo il raggio 2 nel punto P con un certa fase dopo aver percorso una distanza pari a

$$\ell = z \tan \theta. \quad (9.1)$$

Nello stesso punto P arrivano anche tutte le altre onde prodotte in ciascuno dei punti della fenditura. In particolare dal punto centrale della fenditura O parte un'onda che giunge in P dopo essersi propagata lungo 1 e aver percorso una distanza

$$\ell' = z \tan \theta'. \quad (9.2)$$

Quest'ultima interferisce con la prima nel punto P e il risultato consiste in un'onda che può avere un'ampiezza compresa tra la somma delle due ampiezze (interferenza costruttiva) e zero (interferenza distruttiva).

Se la distanza z è molto maggiore dell'ampiezza d della fenditura le onde che partono da un punto estremo e dal centro della fenditura viaggiano praticamente parallele, cioè $\theta \simeq \theta'$. La differenza di cammino percorso dalle due onde $\delta = |\ell - \ell'|$ è molto piccola rispetto a z e si può valutare (guardando la Figura 9.4, che riporta un ingrandimento della Fig. 9.3) come

$$\delta = \frac{d}{2} \sin \theta. \quad (9.3)$$

Infatti, se θ è l'angolo formato tra la normale alla fenditura e la direzione nella quale si trova il punto considerato, riportando la normale alla direzione 1 condotta dal punto A , si vede che il triangolo giallo $AA'O$ è simile al triangolo arancione OHP e quindi l'angolo $\overline{OAA'} = \overline{POH} = \theta$. Essendo $\overline{AO} = \frac{d}{2}$ si trova il risultato che abbiamo scritto sopra.

Se la differenza di cammino δ è pari esattamente a metà della lunghezza d'onda λ dell'onda che si sta propagando, le due onde interferiscono distruttivamente e si cancellano a vicenda producendo nel pun-

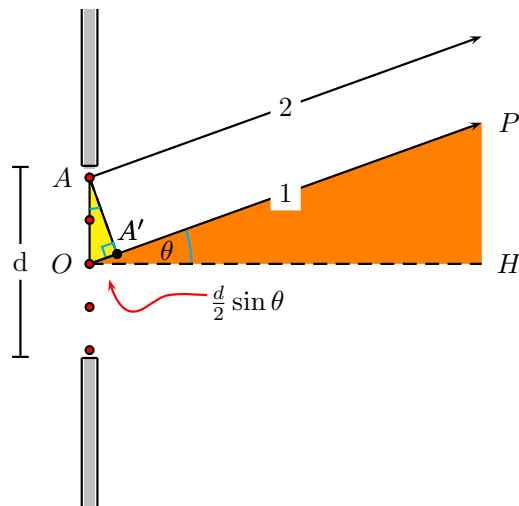


Figura 9.4 Ingrandimento della Figura 9.3 in prossimità della fenditura. Se la distanza z è grande i raggi 1 e 2 sono praticamente paralleli.

to P un'onda di ampiezza nulla. Lo stesso accade se δ è pari a un numero dispari di mezze lunghezze d'onda: $3/2\lambda$, $5/2\lambda$, $7/2\lambda$ e così via. La condizione che determina l'annullarsi dell'onda risultante nel punto P è dunque

$$\delta = (2n + 1) \frac{\lambda}{2} \quad (9.4)$$

che si può riscrivere, sostituendo l'espressione di δ come

$$\frac{d}{2} \sin \theta = (2n + 1) \frac{\lambda}{2}. \quad (9.5)$$

Dunque ci sarà un minimo di ampiezza dell'onda risultante quando $n = 0$ per $\sin \theta = \frac{\lambda}{d}$, un secondo minimo quando $n = 1$ per $\sin \theta = \frac{\lambda}{3d}$, etc.. Se θ è abbastanza piccolo (il che significa che z è grande o che si considerano punti vicini al punto centrale H), chiamando y la distanza tra il punto centrale H e il punto P possiamo scrivere che

$$\sin \theta \simeq \tan \theta = \frac{y}{z} \quad (9.6)$$

e quindi la relazione che stabilisce quando l'interferenza tra le onde considerate è distruttiva si riscrive

$$\sin \theta \simeq \frac{y}{z} = (2n + 1) \frac{\lambda}{d}. \quad (9.7)$$

Naturalmente nel punto P non confluiscono solo le onde che abbiamo considerato, ma tutte quelle prodotte da uno qualunque dei punti della fenditura. Possiamo però sempre trovare una qualunque coppia di punti che distano $d/2$ tra loro all'interno della fenditura per i quali succede esattamente lo stesso. Ad esempio, dal punto rosso compreso tra i punti A e O della Figura 9.4 parte un altro treno di onde nella stessa direzione, così come dal punto immediatamente sotto al punto O . Le onde provenienti da questi due punti fanno esattamente lo stesso di quello che fanno le onde che partono da A e da O . Di queste coppie di punti ce ne sono infinite e per ciascuna valgono le stesse considerazioni fatte sopra: se per certi angoli le onde che viaggiano lungo le direzioni 1 e 2 si cancellano lo fanno anche quelle, a loro parallele, che partono da una qualunque coppia di punti che distano $d/2$ tra loro all'interno della fenditura.

Il risultato è che tutte le onde che partono da ciascuno dei punti della fenditura interferiscono negativamente in un punto che dista y dal centro H , quando

$$y = (2n + 1) \frac{\lambda}{d} z. \quad (9.8)$$

Viceversa, da ciascuna di queste coppie di punti, partiranno onde che interferiranno costruttivamente producendo un'onda di ampiezza massima quando la differenza di cammino è pari a un numero intero di lunghezze d'onda

$$\frac{d}{2} \sin \theta = n\lambda \quad (9.9)$$

per cui i massimi di ampiezza si avranno quando

$$\sin \theta \simeq \frac{y}{z} = 2n \frac{\lambda}{d} \quad (9.10)$$

cioè quando

$$y = 2n \frac{\lambda}{d} z. \quad (9.11)$$

Quando l'ampiezza della fenditura d è molto maggiore della lunghezza d'onda incidente λ , cioè per $d \gg \lambda$ il rapporto λ/d tende a diventare nullo e quindi l'angolo θ al quale si osserva un massimo di ampiezza è tale per cui $\sin \theta \simeq 0$ per ogni n . Analogamente anche i minimi si osservano per angoli tali per cui $\sin \theta = 0$. Questo vuol dire che la luce appare provenire unicamente dalla direzione $\theta = 0$ e quindi non si ha diffrazione.

Al contrario, se λ non è trascurabile rispetto a d , il primo massimo si ha quando $n = 0$ per $y = 0$ come c'era da aspettarsi. Un secondo massimo però si trova a distanza $y = 2z\lambda/d$, poi a $y = 4z\lambda/d$ e così via.

Il risultato è che un'onda piana che si propaga in direzione perpendicolare a uno schermo sul quale è praticata una fenditura produce al di là di questa una caratteristica **figura di diffrazione** che consiste nella comparsa di un'immagine della fenditura in corrispondenza della proiezione del centro della fenditura su una superficie che dista z da essa. L'ampiezza dell'onda diminuisce progressivamente allontanandosi dal centro per diventare nulla a distanza y pari a $z\lambda/d$ (per $n = 0$). Invece di smorzarsi completamente, abbastanza sorprendentemente, l'ampiezza dell'onda risultante al di là dello schermo, comincia invece a risalire per acquisire un nuovo massimo a distanza $y = 2z\lambda/d$.

Se la fenditura ha forma circolare dunque ci aspettiamo di vedere una sua *immagine*⁵ con la stessa forma in corrispondenza del punto centrale H , circondata da un anello a una certa distanza dal bordo della prima immagine, a sua volta circondata da un altro anello e così via. Se invece la forma della fenditura è quella di un taglio, quindi di un rettangolo lungo e stretto, l'immagine avrà la stes-

⁵Possiamo parlare propriamente d'immagine per un'onda luminosa. Nel caso di onde di altro tipo per immagine si deve considerare una specie di *impronta* che potrebbe imprimere l'onda su un materiale opportuno. Ad esempio, nel caso delle onde del mare l'immagine sarebbe costituita dalla forma che assume la spiaggia in seguito all'erosione.

sa forma, ma ai lati dell'immagine si formeranno altre immagini simili (se la larghezza del rettangolo è abbastanza piccola la diffrazione si produce nella direzione della larghezza; in direzione dell'altezza del rettangolo, molto maggiore della lunghezza d'onda, non c'è diffrazione).

Per capire quanto sono distanziate le **frange** chiare e scure è utile misurare l'ampiezza della fenditura in unità di lunghezza d'onda della luce. Quando $d \simeq \lambda$ il rapporto $\lambda/d \simeq 1$ e quindi i minimi d'intensità si hanno per

$$\sin \theta \simeq (2n + 1) \quad (9.12)$$

mentre i massimi si presentano per

$$\sin \theta \simeq 2n \quad (9.13)$$

quindi per $n = 0$, che rappresenta il primo **ordine** delle frange, si ha un massimo per $\theta \simeq 0$ e un minimo per $\sin \theta \simeq 1$ cioè per $\theta \simeq \frac{\pi}{2}$. Si capisce facilmente che n non può essere maggiore di zero, perché altrimenti si avrebbe un minimo per $\sin \theta \simeq 3$ che non è possibile.

Quando $d \ll \lambda$, quindi per esempio per $d \simeq \lambda/10$, quando $n = 0$ otterremmo un minimo per $\sin \theta \simeq 10$! In questo caso non si produce interferenza e l'onda si propaga come un'onda circolare al di là dello schermo. Questo accade quando la lunghezza d'onda è molto più grande delle dimensioni della fenditura. Siamo cioè nelle condizioni analizzate all'inizio del paragrafo: la fenditura si comporta come una singola sorgente di onde sferiche.

Al contrario, se $d \simeq 10\lambda$, il minimo per $n = 0$ si trova per $\sin \theta \simeq 0.1$, cioè per $\theta \simeq 6^\circ$. Oltre al massimo a $\theta = 0$ se ne ottiene uno per $\sin \theta \simeq 0.2$, cioè per $\theta \simeq 12^\circ$. Vedremo quindi un altro minimo in corrispondenza di $n = 1$ quando $\sin \theta \simeq 0.3$ e così via, come si vede dalla Tavola 9.1.

L'intensità di un'onda sferica diminuisce come $1/\ell^2$, dove ℓ è la distanza dalla sorgente. È dunque chiaro che le onde che viaggiano più a lungo avranno un'ampiezza minore di quelle che viaggiano meno. La conseguenza è che man mano che ci si allontana dal centro H l'ampiezza delle onde che interferiscono diminuisce costantemente e quindi an-

n	$\sin \theta_{min}$	$\sin \theta_{max}$	θ_{min}	θ_{max}
0	0.1	0	5.7	0
1	0.3	0.2	17.5	11.5
2	0.5	0.4	30.0	23.6
3	0.7	0.6	44.4	36.9
4	0.9	0.8	64.2	53.1

Tavola 9.1 Quando la fenditura è larga 10 lunghezze d'onda si formano fino a 5 massimi. Il massimo successivo si ha per $\theta = 90^\circ$ quindi non è osservabile (l'onda viaggia parallelamente allo schermo). Naturalmente l'intensità dei massimi diminuisce rapidamente e quindi in pratica ne sono visibili sempre meno.

che l'intensità dei massimi subisce una diminuzione progressiva fino a spegnersi del tutto per angoli molto grandi.

Se si fa un grafico dell'intensità dell'onda risultante (che è proporzionale al quadrato dell'ampiezza dell'onda) in funzione dell'angolo formato con il centro H o della distanza da questo y si trova un andamento simile a quello di Fig. 9.5.

9.3.2 Diffrazione da fenditure multiple

Un altro caso interessante è quello che consiste nel considerare una serie di fenditure parallele molto strette ($d \ll \lambda$): è quello che si chiama un **reticolo di diffrazione**. Nella Figura 9.6 è illustrato un caso con 5 fenditure.

In questo caso ciascuna fenditura si comporta come una sorgente di onde sferiche in fase tra loro. Tra due sorgenti adiacenti c'è una differenza di fase pari a

$$\delta = a \sin \theta \quad (9.14)$$

avendo chiamato a la distanza tra le fenditure. Avremo interferenza costruttiva quando la differenza di

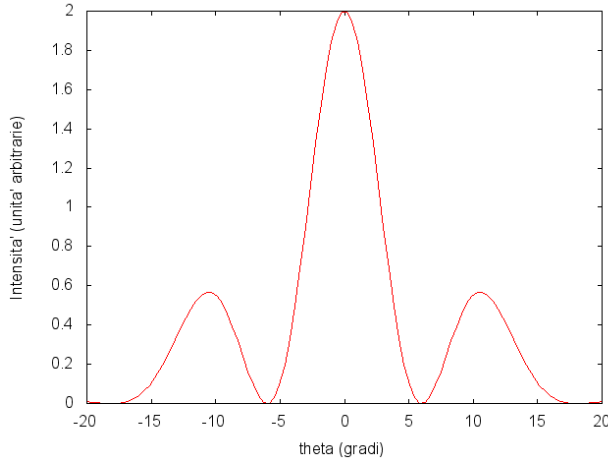


Figura 9.5 Intensità dell'onda che giunge su uno schermo posto al di là di una stretta fenditura, in funzione dell'angolo formato tra il punto dello schermo e la posizione della fenditura. Si vedono il massimo centrale per $\theta = 0$ e i primi massimi secondari a $\theta = \pm 11.5^\circ$. I massimi successivi sono già troppo deboli per potersi apprezzare.

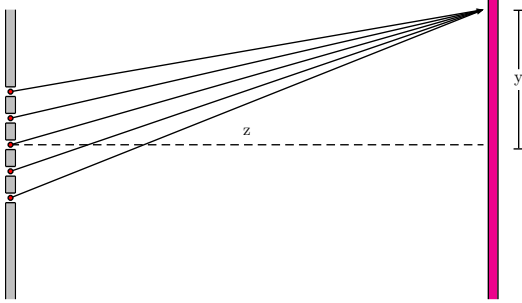


Figura 9.6 Un reticolo di diffrazione con 5 fenditure a distanza z da uno schermo.

fase è pari a un numero intero di lunghezze d'onda λ

$$\delta = a \sin \theta \simeq a \frac{y}{z} = n\lambda \quad (9.15)$$

e cioè quando

$$y = n \frac{\lambda}{a} z. \quad (9.16)$$

Se il reticolo ha una larghezza complessiva d e ha N fenditure equispaziate possiamo scrivere che $a = d/N$, quindi la condizione per ottenere i massimi d'interferenza diventa

$$y_{ret} = n \frac{N\lambda}{d} z. \quad (9.17)$$

Confrontiamo questo risultato con quello ottenuto per la fenditura singola (eq. (9.11))

$$y_{sing} = 2n \frac{\lambda}{d} z. \quad (9.18)$$

Il rapporto tra le distanze comprese tra il massimo centrale e la frangia di ordine $n = 1$ nei due casi vale

$$\frac{y_{ret}}{y_{sing}} = \frac{N}{2}. \quad (9.19)$$

Questo significa che usando un reticolo di diffrazione la separazione tra le frange aumenta di un fattore $N/2$. In altre parole, se vogliamo osservare bene una figura di diffrazione conviene adoperare un reticolo con molte fenditure per unità di lunghezza.

9.4 Qualche esempio pratico

In questo paragrafo descriviamo alcuni esperimenti che si possono condurre in classe o a casa con materiali di facile reperibilità. Gli esperimenti descritti dimostrano la natura ondulatoria del suono e della luce, in quest'ultimo caso risolvendo la disputa di cui abbiamo accennato all'inizio di questo capitolo⁶.

9.4.1 Diffrazione di onde sonore

Il suono è prodotto da onde di compressione dell'aria che giungono fino alle nostre orecchie eccitando

⁶In realtà con la scoperta dell'effetto fotoelettrico tutto verrà rimesso in discussione.

la membrana del timpano, che vibra alla stessa frequenza della sollecitazione prodotta dal moto dell'aria circostante. Queste vibrazioni sono registrate dal cervello che ce le fa percepire come suoni.

Nessuno di noi apparentemente sperimenta quotidianamente effetti di diffrazione sulle onde sonore. Questi dovrebbero comparire come modulazioni d'intensità dei suoni che dipendono dalla posizione assunta rispetto all'apertura attraverso la quale i suoni si propagano al di là di uno schermo. Il motivo è molto semplice: la frequenza ν dei suoni percepibili dall'orecchio umano è compresa tra 20 e 20 000 Hz⁷. La lunghezza d'onda del suono si può ottenere conoscendone la velocità di propagazione in aria che è pari circa a $v = 330$ m/s. Per ricordare la relazione esistente tra frequenza, lunghezza d'onda e velocità, basta ricordare che le frequenze hanno le dimensioni di un tempo alla meno uno, le velocità di lunghezza diviso un tempo e le lunghezze d'onda (evidentemente) di una lunghezza. Pertanto dev'essere

$$\lambda = \frac{v}{\nu}. \quad (9.20)$$

Dunque la lunghezza d'onda dei suoni dev'essere compresa tra

$$\lambda_- = \frac{330}{20} = 16.5\text{m} \quad (9.21)$$

e

$$\lambda_+ = \frac{330}{20\,000} = 16.5\text{mm}. \quad (9.22)$$

Si tratta di un'intervallo molto ampio! I suoni bassi sono quelli con la frequenza più bassa, che dunque hanno la lunghezza d'onda più lunga. Con un semplice *smartphone* potete verificare facilmente che il suono di una grancassa ha una frequenza compresa in un intervallo attorno ai 300 Hz, quindi una lunghezza d'onda tipica del metro. Basta scaricare una qualunque *app* che mostri sullo schermo lo **spettro** del segnale in ingresso sul microfono⁸; soli-

tamente queste *app* impiegano la **Trasformata di Fourier** o FFT (*Fast Fourier Transform*) per ottenerlo. Lo spettro rappresenta l'intensità del segnale in funzione della frequenza.

Al contrario, uno strumento a fiato come un clarinetto produce suoni con frequenza dell'ordine almeno del migliaio di Hertz, per lunghezze d'onda dell'ordine di non più di una trentina di centimetri. Di conseguenza, mentre il suono di una grancassa è in grado di aggirare ostacoli di dimensioni relativamente grandi, il suono dei clarinetti (intorno ai 1500 Hz) non appare diffratto in maniera significativa. Per questa ragione, quando si avvicina una banda musicale e non ci si trova lungo la linea di avanzamento, si sentono prima i suoni dei tamburi e degli strumenti con tonalità più basse e poi quelli del resto della banda: se la banda procede lungo una strada e l'ascoltatore si trova dietro l'angolo in corrispondenza di un incrocio, le onde con frequenza più bassa aggirano gli spigoli e raggiungono l'ascoltatore anche se non si trova lungo la linea di propagazione, grazie alla diffrazione. I suoni con lunghezza d'onda minore della scala tipica delle *fenditure* rappresentate in questo caso dalle strade (di ampiezza pari a qualche metro), invece, non vengono diffratti e si possono udire solo se le onde possono viaggiare in linea retta nella direzione dell'ascoltatore (stiamo trascurando gli effetti delle riflessioni sulle pareti dei caseggiati, in questo caso).

9.4.2 Diffrazione di onde luminose

Un esperimento che mostra chiaramente il fenomeno della diffrazione oggi si può realizzare facilmente con materiali di bassissimo costo⁹. Sono sufficienti un puntatore laser da pochi euro, un CD con almeno una porzione della superficie trasparente e qualche supporto.

La superficie di un CD è costellata di incisioni minutissime disposte a spirale separate da una distanza pari a circa $1.6\text{ }\mu\text{m}$. Illuminando con la luce di un laser una zona trasparente di un CD le trac-

⁷L'intervallo varia molto da individuo a individuo.

⁸Basta cercare sul vostro *store* preferito le parole *spectrum analyzer*. La mia preferita in questo caso è **FFT Spectrum Analyzer** di UberApp.

⁹Era evidentemente molto più difficile nel 1803, quando lo eseguì **Thomas Young** [9].

ce si comportano come le fenditure di un reticolo (il raggio di un CD è molto più grande del raggio dello spot luminoso del laser e le tracce si possono considerare praticamente parallele) e la luce del laser è diffratta. Una parte della luce prosegue nella direzione in cui è puntato il laser, producendo su uno schermo posto al di là della superficie un punto luminoso. Ai lati di questo punto se ne formano altri due, in corrispondenza di un angolo determinato dalla relazione (9.17) a distanza

$$y_{ret} = n \frac{N\lambda}{d} z, \quad (9.23)$$

da quello principale, se la luce si osserva a distanza z dal CD. In questo caso abbiamo $a = d/N = 1.6 \times 10^{-6}$ m, mentre la lunghezza d'onda di un laser rosso è di 635 nm ($\lambda = 635 \times 10^{-9}$ m). Sostituendo i valori si ottiene, per $n = 1$,

$$y_{ret} = \frac{635 \times 10^{-9}}{1.6 \times 10^{-6}} z \simeq 0.40 z. \quad (9.24)$$

Ponendo dunque il CD a circa 30 cm dal muro, su questo si dovrebbero formare almeno tre punti luminosi: uno in corrispondenza del raggio incidente e due equidistanti da questo a una distanza di circa $0.40 \times 30 = 12$ cm. Se il laser è abbastanza potente e la stanza è sufficientemente al buio si potrebbero distinguere anche le frange di ordine $n = 2$. Con un laser verde, la cui lunghezza d'onda è di 520 nm, la distanza tra le frange è di

$$y_{ret} = \frac{520 \times 10^{-9}}{1.6 \times 10^{-6}} 30 \simeq 13 \text{ cm}, \quad (9.25)$$

un po' più ampia. In un DVD le tracce sono più ravvicinate ($a = 0.74 \mu\text{m}$) e quindi l'angolo di diffrazione è maggiore a parità di lunghezza d'onda: le immagini dello spot laser sul muro appaiono più distanziate rispetto a quelle prodotte dal CD. Un disco Blu-Ray ha le tracce che distano solo $0.32 \mu\text{m}$: la distanza tra il massimo centrale e quello di ordine $n = 1$ è quindi quasi doppia rispetto a quella che si può osservare con un DVD.

Per la stessa ragione, osservando la superficie di un CD o un DVD si vedono regioni colorate con

tutti i colori dell'arcobaleno. Il fenomeno della diffrazione, infatti, si può produrre anche attraverso la riflessione e un reticolo funziona anche come dispositivo per **disperdere** la luce. Se s'illumina un reticolo con luce bianca, la luce trasmessa (o riflessa) è scomposta nei diversi colori dell'arcobaleno. In particolare si vedono, ai lati del massimo principale che resta bianco, varie striscie colorate. Questo indica che la luce bianca è una sovrapposizione di onde di lunghezza d'onda diversa, ciascuna delle quali interferisce costruttivamente con le altre ad angoli diversi. Dal momento che l'angolo per il quale si osservano i massimi d'interferenza costruttiva dipende dalla lunghezza d'onda, è evidente che colori diversi corrispondono a lunghezze d'onda diverse. La luce rossa ha una lunghezza d'onda maggiore (e quindi una frequenza minore) rispetto alla luce blu. Conoscendo il passo del reticolo si può dunque misurare la lunghezza d'onda cui corrispondono i diversi colori dello spettro semplicemente misurando l'angolo al quale si vedono i colori.

Questo esperimento permette di stabilire in maniera inequivocabile che **la luce è un'onda**. Non sappiamo ancora in quale mezzo si propaga quest'onda e quindi non sappiamo cosa provoca la sollecitazione che si propaga, ma sappiamo che dev'essere così, perché osserviamo la **diffrazione**.

Un fenomeno analogo si verifica quando vi trovate in auto e osservate le luci delle auto davanti a voi e quelle dei semafori attraverso un parabrezza sporco. Il passaggio dei tergicristalli produce sulla polvere accumulata sul parabrezza una serie di solchi sottilissimi, paralleli e molto vicini: la luce che attraversa questi solchi è diffranta nella direzione perpendicolare alla loro lunghezza e le luci esterne appaiono come *righe tratteggiate* verticali.

9.4.3 Potere risolutivo

Le onde che passano attraverso un foro subiscono inevitabilmente gli effetti della diffrazione. Se il foro è molto grande gli effetti della diffrazione sono piccoli e poco percepibili, ma con un foro molto piccolo gli effetti possono essere importanti.

Una sorgente puntiforme al di là di uno schermo appare come una macchia circolare di raggio pari grosso modo alla distanza tra il centro e il primo minimo. Un foro di diametro d , per esempio, provoca una figura di diffrazione il cui primo minimo si trova a un angolo θ tale che

$$\sin \theta = \frac{\lambda}{d} \quad (9.26)$$

e ponendo uno schermo a distanza z dal foro, l'immagine appare come un cerchio di raggio

$$r = z \tan \theta \simeq z \sin \theta = z \frac{\lambda}{d}. \quad (9.27)$$

Se avessimo due sorgenti puntiformi separate da una distanza angolare α l'una dall'altra, ciascuna delle due produrrebbe sullo schermo un'immagine la cui forma è quella di un cerchio di raggio $r \simeq z\lambda/d$. Perché le due immagini si possano distinguere bene i loro centri devono trovarsi a una distanza pari circa a $2r \simeq 2z\lambda/d$. In effetti, se i due centri si trovano a una distanza l'uno dall'altro pari a un raggio del cerchio, si capisce ugualmente che le sorgenti sono due.

La luce, ad esempio, è un'onda (anche se non abbiamo ancora stabilito di quale natura). Quindi, quando passa attraverso un foro subisce diffrazione. La stessa pupilla dell'occhio umano è un foro attraverso il quale si propagano le onde luminose. Per questa ragione la nostra **acutezza visiva** è limitata: se andate a misurarvi la vista da un ottico vi renderete conto facilmente che, per quanto abbiate la vista buona, riuscirete appena a distinguere le lettere dell'ultima fila nella **tavola ottometrica**. La ragione sta proprio nel fatto che i tratti che formano le lettere sono deformati dalla diffrazione quando la loro luce entra nella vostra pupilla. Al di sotto di certe dimensioni la deformazione subita dai raggi luminosi per effetto della diffrazione diventa comparabile con la dimensione del tratto stesso e le forme delle lettere non sono più distinguibili.

Allo stesso modo, quando di notte osservate i fari di un veicolo molto lontano difficilmente riuscirete a dire se si tratta di una moto o di un'auto, perché le immagini dei fari che si formano sulla retina,

quando sono lontani, si confondono l'una nell'altra per effetto della diffrazione.

Convenzionalmente si stabilisce che il **potere risolutivo** o **risoluzione** r di uno strumento è dato da un numero che è un poco più grande della distanza tra il massimo centrale e il primo minimo

$$r = 1.22 \frac{\lambda}{d} \quad (9.28)$$

dove λ è la lunghezza d'onda della luce osservata attraverso uno strumento con un'apertura di raggio r (**limite di Abbe**)¹⁰.

Anche per questo, per osservare oggetti molto lontani si debbono usare telescopi di grande apertura: la grande apertura (che in certi casi può raggiungere diversi metri) contribuisce a consentire di risolvere dettagli molto fini delle immagini e a raccogliere una maggiore quantità di luce, rendendo gli oggetti deboli più facili da osservare.

9.4.4 Diffrazione di raggi X

La lunghezza d'onda della luce visibile va dai 400 ai 700 nm, circa. Se si invia un fascio di **raggi X**, come quelli che si usano per fare le radiografie, su un reticolo di diffrazione non accade nulla: il fascio prosegue praticamente indisturbato. Questo suggerisce che, se i raggi X sono costituiti di onde come la luce, la loro lunghezza d'onda dev'essere molto più piccola.

Se però s'invia un fascio di raggi X su un cristallo di qualche sostanza, si vede che una parte di questi è riflessa e se se ne misura l'intensità ad angoli diversi si trova una caratteristica figura di diffrazione la cui forma dipende dal tipo di cristallo adoperato. È evidente che questo, da una parte conferma la natura ondulatoria dei raggi X, dall'altra suggerisce che i cristalli siano costituiti di successioni regolari

¹⁰La moderna **fotonica** riesce a superare il limite di Abbe consentendo l'osservazione di strutture molto piccole usando luce visibile, al prezzo di prolungare l'osservazione per un periodo di tempo relativamente lungo. In sostanza quello che si fa non è illuminare l'intero l'oggetto per osservarne la luce diffusa, ma illuminarne un'area molto piccola e registrando l'immagine: il *collage* delle immagini acquisite produce un'immagine complessiva dell'oggetto illuminato.

d'intensità quando $d \sin \theta = m\lambda$.

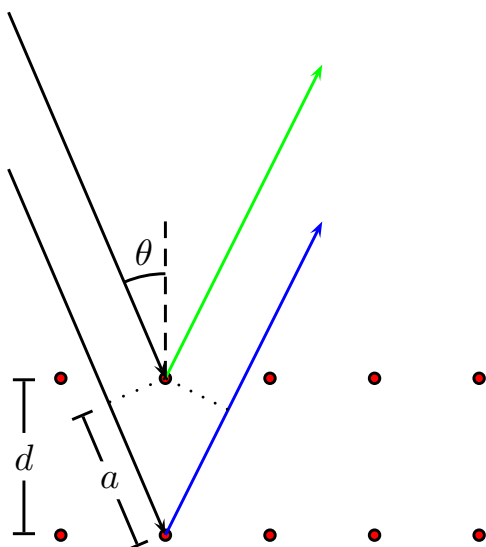


Figura 9.7 Raggi X (in nero) incidono con angolo θ su una serie di centri disposti a strati a distanza d l'uno dall'altro, che li riflettono. I raggi X riflessi sono sfasati tra loro e interferiscono costruttivamente solo a certi angoli. Il raggio verde, ad esempio, percorre una distanza $2a = 2d \sin \theta$ minore rispetto a quello blu. Se questa distanza è pari a un multiplo di una lunghezza d'onda ($2d \sin \theta = m\lambda$) si ha interferenza costruttiva e si osserva un massimo d'intensità.

di centri di diffusione di questo tipo di radiazione. In sostanza, il cristallo deve essere formato da una serie di oggetti disposti in maniera regolare a formare un **reticolo**. Quando il reticolo è investito da un fascio di raggi X, ciascuno di questi oggetti diventa a sua volta una sorgente di raggi X puntiforme. Le diverse onde emesse da questi oggetti interferiscono a certi particolari angoli in modo da formare una figura di diffrazione, la cui forma dipende evidentemente da come sono disposti spazialmente questi centri.

La Figura 9.7 illustra quello che accade. Nel caso semplice di un reticolo regolare si osservano massimi

PREREQUISITI: trasformazioni termodinamiche, primo principio della termodinamica

Se dessimo retta alle leggi fisiche sui gas che discendono dal [primo principio della termodinamica](#) e dall'[equazione di stato dei gas](#), non potremmo spiegare molti fenomeni piuttosto facili da osservare. Ad esempio, non succede mai che l'aria contenuta in una stanza si disponga in modo tale da lasciare del tutto vuota una porzione del volume di quella stanza, anche se, in linea di principio, questo sarebbe possibile: non c'è alcuna legge fisica che lo vieta. Quando si apre una lattina di una bibita gassata, il gas ne viene fuori rapidamente e non succede mai che vi rientra spontaneamente o che resti all'interno della lattina. Non accade mai che un cubetto di ghiaccio si raffreddi ulteriormente, una volta tirato fuori dal freezer, sottraendo calore dall'ambiente circostante, né succede che una tazza di tè si scaldi raffreddando la stanza nella quale lo stiamo bevendo. Di processi come questi ne possiamo trovare a centinaia nella vita quotidiana.

Tutti questi processi hanno in comune il fatto che violano l'**invarianza temporale** intrinseca delle leggi fisiche che li descrivono. La legge che descrive il passaggio di calore tra aria e tè, ad esempio, afferma che la quantità di calore ceduta dal tè dev'essere uguale a quella acquistata dall'aria circostante, ma non dice nulla circa il verso nel quale fluisce il calore che quindi potrebbe passare dall'aria al tè. Perché non lo fa?

In questo capitolo, partendo dal tentativo di realizzare una macchina che funzioni senza alcun intervento esterno e senza somministrazione di una qualsiasi forma di energia, diamo risposta a questa

domanda. È uno di quei casi in cui il tentativo di rispondere a una domanda (perché non si riesce a costruire una macchina che funzioni per sempre senza consumare alcun tipo di carburante?) porta i fisici a scoprire aspetti fondamentali del funzionamento dell'Universo.

10.1 Macchine termiche

Una **macchina** è un dispositivo che compie lavoro. Per esempio, il motore di un'automobile è una macchina perché sposta l'auto e i suoi passeggeri applicando una forza alle ruote; una gru è una macchina perché attraverso l'applicazione di una forza permette di sollevare dei pesi, spostandoli dalla loro posizione originaria; una pompa è una macchina perché sposta un liquido o un gas applicandogli una forza.

Definendo una macchina come qualcosa che compie lavoro, potremmo catalogare tra le macchine anche la stessa Terra, che attraendo verso di sé gli oggetti con la forza di gravità, compie lavoro nei loro confronti. Questo tipo di macchina però non è molto utile: una volta arrivato a terra l'oggetto resta dov'è e non si può più sfruttare la forza di gravità per cambiarne lo stato. Per essere utili le macchine devono essere **cicliche**: devono cioè tornare allo stato iniziale dopo aver compiuto un lavoro. Per esempio, una gru è una macchina utile perché il suo gancio può scendere e salire e quindi si può riportare, una volta fatto un lavoro, nella condizione iniziale. Il motore di un'auto è una macchina utile perché, una volta che il pistone ha compresso il carburante provocandone l'incendio, torna nello stato iniziale ed è pronto per una nuova compressione: se rimanesse

dov'è l'auto non si muoverebbe. Lo stesso vale per una pompa: una volta introdotta l'aria nella ruota di una bicicletta, se non si riporta lo stantuffo nella posizione originale non si riesce a continuare a fare del lavoro.

Le macchine che c'interessa studiare dunque sono le **macchine cicliche**. Una macchina ciclica è qualcosa che si trova in un particolare stato all'inizio, svolge un qualche tipo di lavoro cambiando il proprio stato e ritorna nello stato iniziale. Una macchina reale non torna mai esattamente nello stesso stato di partenza (quanto meno si usura un po', si scalda, etc.). Ma noi possiamo sempre immaginare di poter costruire macchine capaci di farlo, anche se non saremo mai capaci di realizzarle davvero. Eventuali limitazioni di tipo tecnologico hanno poca importanza: prima o poi riusciremo a superarle. Quello che c'interessa fare adesso è capire quali condizioni siamo obbligati a rispettare per realizzare una **macchina ideale**. Una macchina ideale, non reale, è un'approssimazione di una macchina reale che da un lato ci semplifica la vita (possiamo trascurare tutta una serie di effetti che si verificano sulle macchine reali che sarebbero difficili da riprodurre in un modello matematico), dall'altro ci dà indicazioni su quali sono i limiti di questi dispositivi: è chiaro che una macchina reale potrà solo avvicinarsi a una ideale e se qualcosa è impossibile per una macchina ideale a maggior ragione lo sarà per una macchina reale. Da questo punto di vista una macchina ciclica ideale è un qualunque sistema fisico che può compiere un lavoro partendo da uno stato e tornando, passando per una sequenza di stati diversi, a questo. In particolare un gas in un recipiente sottoposto a trasformazioni cicliche si può considerare una macchina tra le più semplici, dal momento che il suo stato dipende esclusivamente dalla temperatura e non dalla posizione dei suoi costituenti. Una tale macchina si chiama **macchina termica** e rappresenta un utile strumento per capire le proprietà delle macchine, perché particolarmente semplice, senza complicazioni di alcun genere.

Il **primo principio della termodinamica** afferma che $\Delta U = \Delta Q - \Delta L$ dove ΔQ è il calore che la macchina assorbe da qualcosa o cede a qualcos'al-

tro e ΔL il lavoro svolto dalla macchina o fatto dall'esterno nei confronti della macchina. Se $\Delta Q > 0$ la macchina **assorbe** calore da qualcosa che si trova all'esterno della macchina (che deve quindi avere una temperatura T_c maggiore di quella della macchina). Se $\Delta Q < 0$ la macchina **cede** calore all'ambiente che la circonda: pertanto la temperatura T_f di questo dev'essere minore di quella della macchina. È evidente che dev'essere $T_c > T_f$ altrimenti la macchina non potrebbe prendere calore dalla **sorgente** a temperatura T_c e cederla alla **sorgente**¹ a temperatura T_f . Prendendo calore dalla sorgente a temperatura T_c deve portare la sua temperatura T a un valore $T + \Delta T$, con $\Delta T < T_c$. Per poter cedere calore alla sorgente a temperatura T_f dev'essere $T + \Delta T > T_f$, cioè dev'essere $T_c > \Delta T > T_f - T$ e quindi che $T_c > T_f - T$. In particolare, visto che T può solo essere positiva, $T_c > T_f$.

Immaginiamo dunque di avere una macchina che compie un lavoro ΔL e che assorbe calore da una sorgente a temperatura T_c , cedendone a una sorgente a temperatura T_f . Chiamiamo ΔQ_c il calore scambiato con la sorgente a temperatura T_c e ΔQ_f quello scambiato con la sorgente a temperatura T_f . Una grandezza fisica di particolare interesse è il **rendimento** η della macchina, definito come

$$\eta = \frac{\Delta L}{\Delta Q_c} \quad (10.1)$$

che ci dice quanto lavoro ΔL riesce a svolgere la macchina a parità di calore assorbito ΔQ_c . Ci piacerebbe, infatti, disporre di una macchina che possa svolgere lavoro senza prendere calore dall'esterno! Se una macchina del genere esistesse potremmo far andare un'automobile senza bruciare carburante, o sollevare pesi con una gru senza un motore o pompare aria in una ruota senza sudare. D'altra parte, a parità di carburante bruciato o di *fatica* fatta, la macchina ha un'efficienza più alta se il lavoro svolto è maggiore. Alla luce del primo principio abbiamo che

¹Chiamiamo **sorgente** un qualunque sistema la cui temperatura sia fissata, sia che assorba che ceda calore.

$$\eta = \frac{\Delta Q - \Delta U}{\Delta Q_c}, \quad (10.2)$$

ma dal momento che la macchina è ciclica $\Delta U = 0$ (l'energia interna U è una **funzione di stato** perciò la variazione di energia interna tra due stati identici non può che essere nulla) e quindi

$$\eta = \frac{\Delta Q}{\Delta Q_c} = \frac{\Delta Q_c - \Delta Q_f}{\Delta Q_c} = 1 - \frac{\Delta Q_f}{\Delta Q_c}. \quad (10.3)$$

10.2 La Macchina di Carnot

La macchina più semplice che possiamo pensare di costruire è fatta con un gas che compie una trasformazione ciclica **reversibile**. Reversibile significa che la macchina passa da uno stato A a uno stato B passando per una serie di stati distinti che si possono ripercorrere all'indietro. Immaginiamo dunque una trasformazione reversibile che porti un gas da uno stato nel quale ha pressione p_A e volume V_A a uno stato in cui ha pressione p_B e volume V_B , per poi tornare nello stato iniziale (p_A, V_A) . Se la trasformazione è reversibile significa che la sequenza di stati per i quali si è passati andando dallo stato A allo stato B si deve poter percorrere al contrario, dallo stato B allo stato A . Perché questo sia possibile è necessario che in ogni istante durante la trasformazione il gas si trovi in uno stato di equilibrio per cui vale $pV = nRT$. Se infatti così non fosse la trasformazione da A a B non sarebbe reversibile: aprendo la lattina di una bibita gassata l'anidride carbonica presente si espande rapidamente uscendo dalla lattina e non è affatto vero che il prodotto della pressione per il volume è uguale a nRT ! Di conseguenza è impossibile riportare il gas fuoriuscito dalla lattina al suo interno facendogli percorrere la trasformazione inversa: non sapremmo nemmeno quale trasformazione fargli subire perché lo stato del gas non è ben definito tra l'istante iniziale e quello finale della trasformazione. Se invece comprimiamo abbastanza lentamente il gas contenuto in un palloncino, la pressione esercitata dal gas sulle pareti del palloncino aumenta progressivamente in modo

tale che $pV = nRT$ in ogni istante. Quindi, diminuendo la pressione sulle pareti, dal momento che il gas continuerebbe a essere in equilibrio si potrebbe tornare nello stato iniziale ripassando per gli stessi stati dai quali si è passati nel corso della compressione. Per questa ragione a volte si dice che le trasformazioni reversibili sono quelle che avvengono lentamente, ma non è sempre vero: lo scioglimento di un ghiacciaio è sicuramente lento, ma non reversibile.

Se la macchina reversibile, qualunque essa sia, fosse in contatto termico con una sola sorgente di temperatura, il gas sarebbe sempre in equilibrio con questa sorgente e perciò avrebbe sempre la stessa temperatura T della sorgente: in altre parole la trasformazione sarebbe **isoterma**, sia all'andata che al ritorno. In definitiva la macchina passerebbe dallo stato A allo stato B percorrendo un ramo d'iperbole nel **piano di Clapeyron**, per poi tornare nello stato A percorrendo all'indietro lo stesso identico ramo d'iperbole. Ricordando che il **lavoro svolto dalla macchina** non è altro che l'area racchiusa all'interno della linea chiusa che rappresenta la trasformazione, si vede subito che in questo caso $\Delta L = 0$, di conseguenza il rendimento $\eta = 0$ e la macchina non sarebbe molto utile!

Un esempio di questo genere di macchina è il seguente: prendiamo un recipiente a forma di cilindro la cui base superiore è mobile e sulla quale si appoggia un peso di massa m . Si riempie il cilindro di gas che si pone in contatto termico con una sorgente a temperatura fissata (basta metterlo in contatto con l'aria avendo cura che l'esperimento abbia una durata non sufficiente a provocare il cambiamento della temperatura di questa). La pressione del gas esercita una forza nei confronti del coperchio che quindi inizia a muoversi verso l'alto. Il moto si arresta quando la forza di gravità cui è soggetto il peso diventa uguale in modulo alla forza esercitata dalla pressione del gas, cioè quando

$$mg = -pA \quad (10.4)$$

se A è l'area di base del cilindro. Il risultato è comunque il compimento di un lavoro $L = mgh$ dove h è lo spostamento verticale del peso, che possiamo

sfruttare per sollevare, appunto, dei pesi. Per utilizzare nuovamente, con un altro peso, questa macchina dobbiamo riportarla nella condizione iniziale. Peccato che per farlo dobbiamo compiere nei confronti del gas un lavoro pari a pAh (dobbiamo cioè esercitare una forza uguale e contraria a quella prodotta dal gas pA per un tratto lungo h). Ma dal momento che $pA = -mg$ i due lavori (quello fatto dal gas e quello fatto da noi per riportare il gas nello stato iniziale) sono identici, ma con segno diverso, per cui il lavoro netto compiuto dal sistema in un ciclo è nullo. In definitiva la macchina è stata inutile: potevamo fare il lavoro di sollevare il peso invece di fare quello di abbassare il pistone. Il risultato sarebbe stato lo stesso.

Perché la macchina sia utile dobbiamo per forza operare tra almeno due temperature T_c e T_f . In questo caso la macchina più semplice che possiamo costruire deve passare da uno stato A a uno stato B percorrendo un'isoterma a temperatura T_c (la macchina in questione potrebbe essere la stessa impiegata nell'esempio precedente allo scopo di sollevare un peso). Per tornare da B ad A siamo però costretti a passare da uno stato a temperatura diversa (quindi, prima di comprimere il gas, dovremmo abbassarne la temperatura in qualche modo). Un'altra isoterma a temperatura $T_f < T_c$ sarebbe rappresentata sul piano di Clapeyron da un ramo d'iperbole più vicino all'asse delle ascisse e non incontrerebbe mai il ramo d'iperbole della prima isoterma. L'unico modo di connettere questi stati tra loro consiste nel fare una prima trasformazione che porta il gas dallo stato B a uno stato C in cui la temperatura passi da T_c a T_f , fargli percorrere un'isoterma fino a uno stato D da cui si raggiunge nuovamente lo stato A (dopo aver compresso il gas dovremmo quindi riscaldarlo). Nei tratti BC e DA non si deve scambiare calore con nessun'altra sorgente se vogliamo averne solo due, perciò queste trasformazioni devono essere **adiabatiche**. Il lavoro fatto da questa macchina è rappresentato dall'area racchiusa nella porzione di piano di Clapeyron delimitata dalle trasformazioni passanti per $ABCD$.

Questo ciclo si chiama **Ciclo di Carnot**² ed è particolarmente importante perché rappresenta il ciclo più semplice possibile con il quale si può produrre lavoro. Tutti gli altri cicli reversibili si possono sempre pensare come somma di cicli di Carnot. Immaginiamo infatti un ciclo qualsiasi, che sul piano di Clapeyron assume una forma qualunque: ogni tratto del ciclo si può sempre pensare decomposto in un tratto di isoterma seguito da un tratto di adiabatica, purché questi tratti siano sufficientemente brevi. Il *perimetro* del ciclo dunque è formato da una serie di isoterme e adiabatiche. L'intera superficie del ciclo, a questo punto, si può *ricoprire* con cicli di Carnot opportuni. Il lavoro compiuto dalla macchina che compie quel ciclo è quindi equivalente al lavoro compiuto da più macchine di Carnot i cui cicli ricoprano completamente quello della nostra macchina. Analizzando dunque le proprietà di un ciclo di Carnot troveremo le proprietà generali di ogni ciclo.

Nel caso di un ciclo reversibile come quello di Carnot, il calore è scambiato solo durante le trasformazioni isoterme (nelle adiabatiche non c'è scambio di calore). Evidentemente il calore è assorbito durante l'espansione isoterma che avviene a temperatura T_c e ceduto nel corso della compressione isoterma a temperatura T_f . Nelle isoterme la temperatura non varia perciò $\Delta U = 0$, quindi $\Delta Q = \Delta L$. Il lavoro fatto durante l'espansione a temperatura T_c vale

$$\Delta L_c = nRT_c \log \frac{V_B}{V_A} = \Delta Q_c \quad (10.5)$$

mentre quello fatto nel corso della compressione a temperatura T_f (che è negativo), vale

$$\Delta L_f = -nRT_f \log \frac{V_C}{V_D} = -\Delta Q_f \quad (10.6)$$

Pertanto possiamo scrivere il rendimento come

$$\eta = 1 - \frac{T_f \log \frac{V_C}{V_D}}{T_c \log \frac{V_B}{V_A}}. \quad (10.7)$$

²Dal nome del fisico **Sadi Carnot** (1796–1832) che per primo ne illustrò le proprietà.

I rapporti tra i volumi sono determinati dal fatto che V_B e V_C , così come V_D e V_A sono punti raggiunti nel corso di un'espansione e una compressione adiabatica, per cui per essi vale l'equazione

$$TV^{\gamma-1} = \text{const} \quad (10.8)$$

quindi possiamo scrivere che

$$T_A V_A^{\gamma-1} = T_D V_D^{\gamma-1} \quad (10.9)$$

e che

$$T_B V_B^{\gamma-1} = T_C V_C^{\gamma-1}. \quad (10.10)$$

Dividendo membro a membro queste equazioni si trova che

$$\frac{T_B}{T_A} \left(\frac{V_B}{V_A} \right)^{\gamma-1} = \frac{T_C}{T_D} \left(\frac{V_C}{V_D} \right)^{\gamma-1} \quad (10.11)$$

ed essendo $T_A = T_B$ e $T_C = T_D$ si ottiene che

$$\frac{V_B}{V_A} = \frac{V_C}{V_D}. \quad (10.12)$$

Scopriamo così che gli argomenti dei logaritmi a numeratore e denominatore nella formula che dà il rendimento di questa macchina sono uguali. Quindi lo sono anche i loro logaritmi e il loro rapporto fa semplicemente 1. Di conseguenza il rendimento si scrive come

$$\eta \equiv 1 - \frac{T_f}{T_c}. \quad (10.13)$$

Si vede subito che, dovendo essere $T_c > T_f$, η è sempre minore di 1 e maggiore di zero. Questo significa che non si può fare una macchina che trasformi in lavoro tutto il calore assorbito, per la quale $\eta = 1$. Questo sarebbe possibile solo se T_c tendesse all'infinito: in quest'ultimo caso il rapporto T_f/T_c tenderebbe a zero e η a uno. Il prezzo da pagare è che occorre mettere in contatto la macchina con una sorgente a temperatura infinita, il che evidentemente non è possibile. A ben guardare anche nel caso in cui $T_f = 0$ il rendimento sarebbe uguale a uno. Ma anche questo è impossibile. Infatti, supponiamo di riuscire a realizzare una sorgente a temperatura

$T = 0$ K. Quando la mettiamo in contatto con la nostra macchina, questa deve cederle calore e questo non può che innalzare la sua temperatura, anche se di pochissimo. Se però $T_f \neq 0$ il rendimento è $\eta < 1$ e da questo momento in poi la macchina funzionerà in questo regime. Nel caso in cui $T_f = T_c$, inoltre, si vede subito che $\eta = 0$, ma questo lo sappiamo già perché in questo caso la macchina non fa lavoro (utile).

Una macchina di Carnot efficiente deve lavorare tra due temperature molto diverse tra loro, in modo tale che il rapporto T_f/T_c sia il più piccolo possibile, ma in ogni caso questo rapporto sarà strettamente maggiore di zero e così il rendimento sarà strettamente minore di uno.

Tutte le macchine reversibili, per quanto detto sopra, si possono sempre rappresentare come macchine che compiono una serie di cicli di Carnot, perciò il loro rendimento si scrive sempre come

$$\eta = 1 - \frac{T_f}{T_c}. \quad (10.14)$$

Le macchine **irreversibili**, invece, non possono che fare meno lavoro di quelle reversibili: per loro dunque vale la relazione

$$\eta \leq 1 - \frac{T_f}{T_c} \quad (10.15)$$

È utile osservare che questa relazione vale per ogni macchina. Ogni macchina, anche se non è una macchina termica, deve rispettare il primo principio della termodinamica per cui nel caso in cui $\Delta Q = 0$, può compiere lavoro solo se cambia la sua energia interna. Ma se la macchina fosse perfetta e la sua temperatura non cambiasse nel corso del ciclo che produce lavoro, la variazione di energia interna non potrebbe che derivare da una variazione dell'energia potenziale dei suoi costituenti: questo significa che l'organizzazione interna della macchina cambierebbe nel tempo. Prima o poi, dunque, i costituenti di questa macchina occuperanno posizioni relative diverse da quelle iniziali e di conseguenza la macchina non potrà mai tornare nelle condizioni di partenza (a meno che non si fornisca calore o si faccia nei

confronti della macchina un lavoro per riportare i suoi costituenti nelle posizioni iniziali).

Questa scoperta, avvenuta alla fine del XVIII secolo, decretò la fine delle ricerche che avevano lo scopo di realizzare il **moto perpetuo**, cioè di un motore che potesse funzionare per sempre senza essere ricaricato in qualche modo e senza l'apporto di carburanti o di altre fonti esterne.

10.3 Entropia

Il rendimento di una qualunque macchina è, per definizione, il rapporto tra il lavoro ΔL che riesce a fare e la quantità di calore ΔQ_c che si deve fornire dall'esterno per farla funzionare:

$$\eta = \frac{\Delta L}{\Delta Q_c} = \frac{\Delta Q_c - \Delta Q_f}{\Delta Q_c} = 1 - \frac{\Delta Q_f}{\Delta Q_c}. \quad (10.16)$$

Per ogni macchina vale anche che

$$\eta \leq 1 - \frac{T_f}{T_c} \quad (10.17)$$

e quindi si deve avere che

$$1 - \frac{\Delta Q_f}{\Delta Q_c} \leq 1 - \frac{T_f}{T_c}, \quad (10.18)$$

da cui segue che

$$\frac{\Delta Q_f}{\Delta Q_c} \geq \frac{T_f}{T_c}. \quad (10.19)$$

Possiamo anche scrivere questa relazione nella forma

$$\frac{\Delta Q_f}{T_f} \geq \frac{\Delta Q_c}{T_c}. \quad (10.20)$$

10.4 Il secondo principio della termodinamica

Per le macchine reversibili evidentemente vale la relazione di uguaglianza

$$\frac{\Delta Q_f}{T_f} = \frac{\Delta Q_c}{T_c}, \quad (10.21)$$

che si può tradurre dicendo che il rapporto tra il calore scambiato in una trasformazione, diviso per la temperatura alla quale avviene la trasformazione è costante (evidentemente, qualunque sia la trasformazione, si può sempre dividere in un numero arbitrario di passi abbastanza piccoli da poter considerare costante la temperatura in ciascuno di essi). Al rapporto $\Delta S = \Delta Q/T$ si dà il nome di **entropia** o, meglio, **variazione di entropia**, intendendo che il rapporto in questione si può definire per una trasformazione che porta una macchina da uno stato A a uno stato B durante la quale la macchina assorbe o cede una quantità di calore ΔQ . Se la macchina non cambia stato la sua variazione di entropia è evidentemente nulla, così come è nulla la variazione di entropia nel corso di una trasformazione adiabatica (essendo $\Delta Q = 0$). Per un ciclo reversibile la somma delle variazioni di entropia è nulla. Infatti, se il rapporto $\Delta Q/T$ è costante allora

$$\frac{\Delta Q_c}{T_c} - \frac{\Delta Q_f}{T_f} = 0. \quad (10.22)$$

Questo risultato si può scrivere nella forma

$$\sum_i \frac{\Delta Q_i}{T_i} = 0, \quad (10.23)$$

dove i segni di ΔQ_i si prendono secondo le usuali convenzioni (positivo se il calore entra nella macchina, negativo se esce). Se la variazione di entropia in una trasformazione ciclica è nulla, come sempre questo significa che l'entropia è una **funzione di stato**: dipende cioè solo dallo stato della macchina e non dal particolare modo in cui tale stato è stato raggiunto. Di conseguenza la variazione di entropia dipende solo dagli stati iniziale e finale e non dalla particolare trasformazione subita dalla macchina.

Per le macchine cicliche irreversibili, dal momento che

$$\frac{\Delta Q_f}{\Delta Q_c} \geq \frac{T_f}{T_c}, \quad (10.24)$$

vale che, in un ciclo,

$$\frac{\Delta Q_f}{T_f} \geq \frac{\Delta Q_c}{T_c}. \quad (10.25)$$

Possiamo riscrivere la relazione come

$$\frac{\Delta Q_c}{T_c} - \frac{\Delta Q_f}{T_f} \leq 0, \quad (10.26)$$

cioè che

$$\sum_i \frac{\Delta Q_i}{T_i} \leq 0. \quad (10.27)$$

In questo caso, evidentemente, il rapporto $\Delta Q/T$ **non è una funzione di stato** (in una trasformazione ciclica la somma di questi rapporti non è nulla, il che significa che la somma dipende dalla particolare trasformazione fatta).

Dal momento che per una trasformazione ciclica reversibile $\Delta S = 0$ possiamo sostituire a zero il simbolo ΔS e scrivere che

$$\sum_i \frac{\Delta Q_i}{T_i} \leq \Delta S, \quad (10.28)$$

cioè che

$$\Delta S \geq \sum_i \frac{\Delta Q_i}{T_i}. \quad (10.29)$$

Se però ΔS è una funzione di stato, la stessa relazione vale qualunque sia la trasformazione. Se quindi si prende una qualunque trasformazione che porta la macchina da uno stato A a uno stato $B \neq A$, possiamo scrivere che

$$\Delta S_{A \rightarrow B} \geq \sum_i \frac{\Delta Q_i}{T_i} \bigg|_{A \rightarrow B}, \quad (10.30)$$

intendendo che la somma va estesa a tutte le porzioni di trasformazione che portano il sistema da A a B . La somma sulla destra prende il nome di **integrale di Clausius**, dal nome di **Rudolf Clausius** che **coniò** il termine di entropia per il rapporto $\Delta Q/T$.

La conseguenza è che in un sistema **isolato**, che non scambia calore con l'esterno, $\Delta Q_i = 0$ e quindi $\Delta S \geq 0$, cioè la variazione di entropia non è mai negativa o, il che è lo stesso, l'entropia non può che aumentare. Questo risultato prende il nome di **secondo principio della termodinamica** ed è particolarmente importante, non soltanto per le implicazioni sulle possibilità di realizzare macchine efficienti.

In primo luogo quella appena enunciata è una legge fisica fondamentale che deriva da una **teoria**: la stessa seconda legge di Newton per cui $\mathbf{F} = m\mathbf{a}$ non è altrettanto fondamentale, perché in fondo non fa altro che tradurre in termini formali un'osservazione sperimentale. Il secondo principio della termodinamica invece, pur prendendo le mosse da osservazioni sperimentali (e non potrebbe essere altrimenti), deriva dall'analisi delle proprietà matematiche del rapporto $\Delta L/\Delta Q_c$. Il risultato, cioè, non è direttamente deducibile dall'esperienza. L'altro motivo per cui questo risultato assume un'importanza particolare è che stabilisce la direzione nella quale fluisce il tempo: quella in cui l'entropia dell'Universo aumenta! Vale la pena fare due precisazioni.

1. La variazione di entropia per portare un qualunque sistema da uno stato A a uno stato $B \neq A$ si può calcolare scegliendo una o più trasformazioni reversibili qualunque che portino il sistema dallo stato A allo stato B , indipendentemente dalla trasformazione che il sistema ha effettivamente subito e indipendentemente dal fatto che la trasformazione sia o meno reversibile.
2. La variazione di entropia di un sistema può essere sia positiva che negativa (in un sistema l'entropia può quindi aumentare o diminuire), ma non può essere negativa (l'entropia non può diminuire) se il sistema è isolato.

Il sistema isolato per eccellenza è l'**Universo**, il quale non può scambiare calore o fare lavoro con nient'altro, perciò possiamo affermare che **l'entropia dell'Universo può solo aumentare**. Nei paragrafi che seguono calcoliamo la variazione di entropia di alcuni sistemi notevoli, per prendere confidenza con questa grandezza fisica, le cui dimensioni sono quelle di un'energia diviso una temperatura e che perciò si misura in J/K.

Prima però osserviamo che lo studio del rendimento dei motori ci ha portato a una conclusione estremamente importante: l'entropia dell'intero Universo non può mai diminuire. Questo fissa, in un certo senso, la direzione nella quale scorre il tempo:

i processi si possono svolgere soltanto se portano a un aumento dell'entropia dell'Universo. Se la variazione di entropia di un determinato processo porta l'entropia dell'Universo a diminuire, quel processo è vietato.

Trattandosi della conseguenza di una teoria, il secondo principio va messo alla prova eseguendo una serie di esperimenti i cui risultati si devono poter predire in base a questa teoria. Vediamo quindi una serie di esperimenti che si possono condurre per verificare quanto sopra.

10.4.1 Passaggi di calore a volume costante

Quando un sistema è scaldato, la sua temperatura passa dal valore iniziale T_i a un valore finale $T_f > T_i$ in seguito all'assorbimento di una certa quantità di calore ΔQ che si può scrivere come $\Delta Q = C\Delta T$ dove C è la capacità termica del sistema. Se il sistema è un solido o un liquido $C = mc$ con c pari al calore specifico della sostanza di cui è formato. Nel caso dei gas perfetti $C = nc_V = \frac{3}{2}nR$. La variazione di entropia subita dal sistema la possiamo perciò scrivere come

$$\Delta S = \sum_i \frac{\Delta Q}{T} = \sum_i \frac{C\Delta T}{T}. \quad (10.31)$$

Dal momento che durante la trasformazione la temperatura cambia, non possiamo calcolare banalmente la somma a destra dell'equazione. L'analisi matematica permette di eseguire il calcolo che, come risultato, dà

$$\Delta S = C \log \frac{T_f}{T_i}. \quad (10.32)$$

Poiché $T_f > T_i$ il logaritmo è positivo, così come la variazione di entropia (sia che si tratti di un solido, di un liquido o di un gas).

Se invece di riscaldare il sistema lo raffreddiamo quello che succede è che non cambia l'espressione della variazione di entropia, ma il suo segno sì. Quando si raffredda un sistema, dunque, la sua

entropia diminuisce. Questo non è in contraddizione con il fatto che l'entropia complessiva dell'Universo deve aumentare: se l'entropia di una parte dell'Universo (quella costituita dal sistema in esame) diminuisce, l'entropia del resto dell'Universo deve aumentare in modo tale che complessivamente l'entropia dell'Universo aumenti (o al più resti costante).

Un caso particolarmente semplice si ha quando si provoca la solidificazione di una quantità m di acqua che si trasforma in ghiaccio. Perché possa avvenire questo fenomeno è necessario sottrarre all'acqua una quantità di calore $\Delta Q = \lambda m$, dove λ è il **calore latente** di solidificazione. Durante il processo la temperatura T resta costante e pari a $T = 273$ K circa. La variazione di entropia dell'acqua dunque vale

$$\Delta S = -\frac{\lambda m}{T} \quad (10.33)$$

dove il segno $-$ indica che all'acqua si sta sottraendo calore. Quando il ghiaccio si scioglie, invece, acquista calore perciò la sua variazione di entropia sarà la stessa, ma positiva:

$$\Delta S = \frac{\lambda m}{T}. \quad (10.34)$$

Nei processi che abbiamo descritto la variazione di entropia del sistema che consideriamo può essere sia positiva che negativa, perché il sistema non è mai isolato (altrimenti non potrebbe scambiare calore con l'esterno). D'altra parte il ghiaccio si scioglie solo se è in contatto termico con qualcosa a temperatura maggiore, come l'aria, ad esempio. La variazione di entropia dell'aria sarà $\Delta Q/T_{aria}$, ma dal momento che $T_{aria} > T$ e che resta praticamente costante, mentre $\Delta Q = \lambda m$ dev'essere pari a quella necessaria per sciogliere il ghiaccio, si vede subito che

$$\begin{aligned} \Delta S_U &= \Delta S_{ghiaccio} + \Delta S_{aria} \\ &= \lambda m \left(-\frac{1}{T} + \frac{1}{T_{aria}} \right) > 0. \end{aligned} \quad (10.35)$$

Consideriamo ora un sistema completamente isolato, che non scambia calore con altri sistemi né compie lavoro, come una coppia di oggetti a temperatura diversa T_c e T_f con $T_c > T_f$. Se posti a contatto, i due oggetti raggiungono la stessa temperatura di equilibrio T . Se, per semplicità, i due oggetti sono uguali e fatti dello stesso materiale la loro capacità termica è la stessa e

$$T = \frac{T_c + T_f}{2}. \quad (10.36)$$

La variazione di entropia dell'oggetto inizialmente più freddo è positiva e vale

$$\Delta S_f = C \log \frac{T}{T_f} \quad (10.37)$$

($T > T_f$, quindi il logaritmo è positivo) mentre quella dell'oggetto inizialmente più caldo è negativa, perché $T < T_c$ e vale

$$\Delta S_c = C \log \frac{T}{T_c}. \quad (10.38)$$

La variazione di entropia totale ΔS_U (dell'Universo, diremmo) è la somma delle due variazioni di entropia:

$$\begin{aligned} \Delta S_U &= \Delta S_c + \Delta S_f = C \left(\log \frac{T}{T_c} + \log \frac{T}{T_f} \right) \\ &= C \log \frac{T^2}{T_c T_f} \end{aligned} \quad (10.39)$$

dove abbiamo sfruttato la proprietà dei logaritmi per cui la somma dei logaritmi è il logaritmo del prodotto. Sostituendo nell'espressione sopra trovata quella di T abbiamo che

$$\Delta S_U = C \log \frac{(T_c + T_f)^2}{4T_c T_f}. \quad (10.40)$$

Se poniamo $T_c = T_f + \delta$, l'argomento del logaritmo si scrive come

$$\frac{(2T_f + \delta)^2}{4(T_f^2 + \delta T_f)} = \frac{4T_f^2 + 4T_f \delta + \delta^2}{4T_f^2 + 4T_f \delta} \quad (10.41)$$

e si vede subito che il numeratore è maggiore del denominatore, perciò l'argomento del logaritmo è maggiore di uno e la variazione di entropia complessiva è positiva: $\Delta S_U \geq 0$, come ci aspettavamo³. L'entropia diminuirebbe solo se T_c aumentasse e T_f diminuisse e cioè se il corpo freddo diventasse più freddo e quello caldo più caldo. Vale la pena notare che un caso come questo non sarebbe vietato dalle altre leggi fisiche, incluso il primo principio della termodinamica. Nulla vieta, infatti, che il calore si propaghi dai corpi freddi a quelli caldi rendendo quelli freddi ancor più freddi! In effetti la propagazione del calore da un corpo all'altro avviene perché le parti di cui è composto un corpo caldo si muovono più rapidamente di quelle di cui è composto un corpo freddo. Se le parti del corpo caldo urtano quelle del corpo freddo possono trasferire energia cinetica a queste ultime, riscaldandole. Ma potrebbe anche succedere che nell'urto sia la particella del corpo caldo a sottrarre energia da quella del corpo freddo, provocandone un abbassamento della temperatura. Questo processo non è vietato da alcuna legge fisica, purché l'energia totale si conservi! L'unica legge che impone che i corpi freddi si scaldino a contatto con quelli caldi è quella che abbiamo trovato in questo capitolo e cioè che in ogni trasformazione l'entropia dell'Universo deve necessariamente aumentare o, al più, restare costante.

10.5 L'espansione irreversibile di un gas

Abbiamo già osservato che, aprendo la lattina di una bibita gassata, il gas contenuto nella lattina si espande rapidamente uscendo dalla lattina e non succede mai che dell'aria penetri nella lattina aumentando la quantità di gas presente all'interno. In effetti, quando si apre una lattina, il gas passa da uno stato di equilibrio in cui pressione e volume sono tali per cui vale

³Vale la pena osservare che per $\delta = 0$ l'argomento del logaritmo vale 1 e quindi la variazione di entropia è nulla.

$$pV = nRT \quad (10.42)$$

a uno stato in cui la pressione diventa improvvisamente pari a quella atmosferica (dunque inferiore a p) e il volume a disposizione aumenta improvvisamente. Nel volgere di alcuni secondi il gas, fuoriuscendo dalla lattina, raggiunge l'equilibrio termico con l'aria circostante. Supponiamo di realizzare l'esperimento con una lattina tirata fuori da un frigorifero a $T = 4^\circ \text{C}$, pari a $T \simeq 277 \text{ K}$. In una tipica lattina da 33 cl ci sono circa⁴ $n = 0.05$ moli di CO_2 , perciò la pressione del gas è pari a

$$p = \frac{nRT}{V} \simeq \frac{0.05 \times 8.314 \times 277}{0.33 \times 10^{-3}} \simeq 349 \text{ kPa}. \quad (10.43)$$

Nello stato finale il gas ha una pressione p_f pari a quella atmosferica $p_f \simeq 101 \text{ kPa}$ e una temperatura pari a quella ambiente $T \simeq 20^\circ \text{C} = 293 \text{ K}$. L'espansione e il riscaldamento del gas non sono affatto trasformazioni reversibili e negli stati intermedi tra quello iniziale e quello finale l'equazione di stato dei gas non è valida. Nonostante ciò possiamo calcolare la variazione di entropia del gas trovando una qualunque trasformazione che porti il gas dallo stato in cui la pressione vale $p_i = 349 \text{ kPa}$ e la temperatura vale $T_i = 277 \text{ K}$, a uno stato in cui pressione e temperatura valgono $p_f = 101 \text{ kPa}$ e $T_f = 293 \text{ K}$, rispettivamente. Un gas perfetto in queste condizioni avrebbe un volume

$$V_f = \frac{nRT}{p} \simeq 1.21 \times 10^{-3} \text{ m}^3. \quad (10.44)$$

Ci sono infiniti modi per passare dallo stato iniziale a quello finale con trasformazioni reversibili: uno di questi consiste nel far espandere il gas a pressione costante da V_i a V_f per poi portare la sua pressione a p_f lasciando costante il volume (il gas compie

⁴La stima è stata fatta assumendo che il gas abbia a disposizione tutto il volume all'interno della lattina; evidentemente non è così. Inoltre, nelle condizioni in cui si trova il gas, quest'ultimo non si può considerare ideale, quindi l'equazione di stato non vale esattamente. Per i nostri fini, tuttavia, la quantità di gas presente è un numero arbitrario.

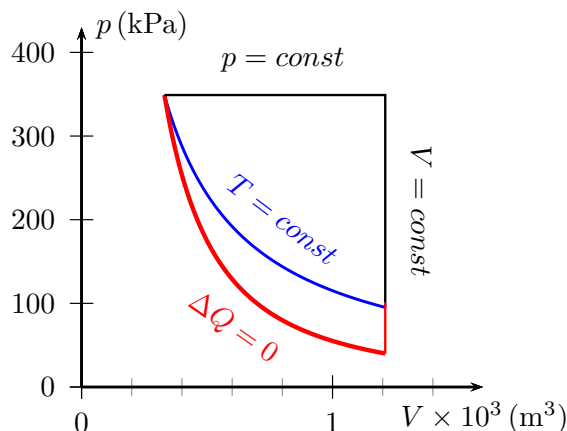


Figura 10.1 Tre possibili trasformazioni che portano un gas da uno stato iniziale in cui $p_i = 349 \text{ kPa}$ e $V_i = 33 \text{ cl}$, a uno stato in cui $p_f = 101 \text{ kPa}$ e $V_f = 121 \text{ cl}$: un'isobara seguita da un'isocora (in nero), un'isoterma seguita da un'isocora (in blu) e un'adiabatica seguita da un'isocora (in rosso). Il tratto di isocora blu è poco visibile perché brevissimo.

dunque una trasformazione isobara seguita da un'isocora, come quella in nero della Figura 10.1); in alternativa si potrebbe far compiere al gas una trasformazione isoterma (in blu nella stessa figura) che lo porterebbe a una pressione di poco diversa da quella finale, seguita da una trasformazione isocora, oppure si potrebbe fargli compiere una trasformazione adiabatica (in rosso) per portarla al volume finale, seguita ancora una volta da un'isocora per portare il gas alla pressione di equilibrio.

Calcoliamo la variazione di entropia nelle diverse trasformazioni, iniziando da quella isobara che porta il gas dallo stato (p_i, V_i) allo stato (p_i, V_f) . Durante questa trasformazione la temperatura cambia quindi dobbiamo dividere la trasformazione in tratti sufficientemente brevi da poterci considerare la temperatura costante, in ciascuno dei quali possiamo scrivere che

$$\Delta S_i = \frac{\Delta Q_i}{T_i} = \frac{nc_p \Delta T_i}{T_i}, \quad (10.45)$$

dove $c_p = c_v + R$ è il calore specifico a pressione costante e $c_v = \frac{3}{2}R$ quello a volume costante. Come nel caso dell'equazione (10.32) la somma $\sum_i \Delta S_i$ si scrive come

$$nc_p \log \frac{T}{T_i} \quad (10.46)$$

in cui T rappresenta la temperatura del gas nello stato (p_i, V_f) che si ottiene dall'equazione di stato per cui

$$T = \frac{p_i V_f}{nR}. \quad (10.47)$$

In definitiva abbiamo che la variazione di entropia ΔS_p del gas nel corso dell'espansione isobara è

$$\Delta S_p = nc_p \log \frac{p_i V_f}{nRT_i}, \quad (10.48)$$

che per l'equazione di stato dei gas che assicura che $nRT_i = p_i V_i$, si può riscrivere come

$$\Delta S_p = nc_p \log \frac{p_i V_f}{p_i V_i} = nc_p \log \frac{V_f}{V_i}, \quad (10.49)$$

In maniera del tutto analoga, nel corso dell'isocora si varia l'entropia ΔS_V del gas di

$$\Delta S_V = nc_v \log \frac{T_f}{T} = nc_v \log \frac{p_f}{p_i}. \quad (10.50)$$

Sommando i due contributi si ottiene

$$\Delta S = \Delta S_p + \Delta S_V = \frac{5}{2}nR \log \frac{V_f}{V_i} + \frac{3}{2}nR \log \frac{p_f}{p_i}. \quad (10.51)$$

Eseguiamo ora lo stesso conto assumendo che il gas compia prima la trasformazione isoterma e poi l'isocora. Durante l'isoterma la temperatura resta costante perciò la variazione di entropia ΔS_T si scrive

$$\Delta S_T = \frac{\Delta Q_T}{T} \quad (10.52)$$

dove, dal momento che in un'isoterma $\Delta U = 0$, la quantità di calore scambiata $\Delta Q_T = \Delta L_T$, dove con ΔL_T abbiamo indicato il lavoro fatto nel corso della trasformazione che possiamo calcolare come l'area sottesa dalla curva dell'isoterma sul piano (p, V) : dividendo l'intervallo tra V_i e V_f in piccoli intervalli di ampiezza ΔV_i , in ciascuno possiamo considerare la pressione p_i costante e scrivere che

$$\Delta L = \sum_i \Delta L_i = \sum_i p_i \Delta V_i, \quad (10.53)$$

ma vale sempre che $p_i = nRT/V_i$ e dunque

$$\Delta L = \sum_i nRT \frac{\Delta V_i}{V_i}, \quad (10.54)$$

uguale, per la solita ragione, a

$$\Delta L = nRT \log \frac{V_f}{V_i}. \quad (10.55)$$

Di conseguenza

$$\Delta S_T = nR \log \frac{V_f}{V_i}. \quad (10.56)$$

Quando si percorre il tratto di isocora possiamo usare il risultato ottenuto sopra, all'equazione (10.50)

$$\Delta S_V = \frac{3}{2}nR \log \frac{T_f}{T} = \frac{3}{2}nR \log \frac{p_f}{p_i}. \quad (10.57)$$

in cui, al posto di p_i c'è la pressione p che il gas avrebbe avuto avendo raggiunto il punto finale della prima isoterma per la quale $T = T_i$ e $V = V_f$, quindi

$$p = \frac{nRT_i}{V_f} = p_i \frac{V_i}{V_f}. \quad (10.58)$$

Sostituendo si trova

$$\Delta S_V = \frac{3}{2}nR \log \frac{p_f V_f}{p_i V_i}. \quad (10.59)$$

La somma dei due contributi dà

$$\Delta S = \Delta S_T + \Delta S_V = nR \log \frac{V_f}{V_i} + \frac{3}{2}nR \log \frac{p_f V_f}{p_i V_i}, \quad (10.60)$$

che, grazie alle proprietà dei logaritmi, diventa

$$\Delta S = nR \left(1 + \frac{3}{2} \right) \log \frac{V_f}{V_i} + \frac{3}{2} nR \log \frac{p_f}{p_i}, \quad (10.61)$$

che è esattamente quanto abbiamo ottenuto all'equazione (10.51), come del resto dovevamo aspettarci sapendo che l'entropia è una **funzione di stato** la cui variazione dipende solo dagli stati iniziale e finale e non dalla particolare trasformazione fatta. È evidente, a questo punto, che lo stesso risultato dobbiamo ottenerlo nel caso della terza serie di trasformazioni: un'adiabatica seguita da un'isocora. In questo caso, poiché durante l'adiabatica non si scambia calore, la variazione di entropia $\Delta S_A = 0$ e basta calcolare solo quella derivante dalla trasformazione isocora, che è sempre

$$\Delta S_V = \frac{3}{2} nR \log \frac{p_f}{p_i}. \quad (10.62)$$

Qui la pressione iniziale p è quella che il gas avrebbe raggiunto se avesse subito la trasformazione adiabatica partendo dallo stato (p_i, V_i) . Dal momento che durante un'espansione adiabatica il prodotto pV^γ , con $\gamma = c_p/c_v = 5/3$ resta costante, abbiamo che

$$pV_f^\gamma = p_i V_i^\gamma \quad (10.63)$$

da cui si ottiene

$$p = p_i \left(\frac{V_i}{V_f} \right)^\gamma \quad (10.64)$$

che, sostituita al posto di p_i dell'equazione che ci dà ΔS_V conduce al risultato

$$\Delta S = \Delta S_V = \frac{3}{2} nR \log \frac{p_f}{p_i} \left(\frac{V_f}{V_i} \right)^{\frac{5}{3}}. \quad (10.65)$$

Grazie alle proprietà dei logaritmi possiamo riscrivere quest'equazione come

$$\Delta S = \frac{3}{2} nR \log \frac{p_f}{p_i} + \frac{3}{2} nR \log \left(\frac{V_f}{V_i} \right)^{\frac{5}{3}} \quad (10.66)$$

e quindi, scrivendo che $\log a^b = b \log a$, come

$$\Delta S = \frac{3}{2} nR \log \frac{p_f}{p_i} + \frac{5}{2} nR \log \frac{V_f}{V_i} \quad (10.67)$$

che è ancora uguale all'equazione (10.51).

In generale, quando si deve valutare una variazione di entropia, proprio allo scopo di semplificare i calcoli, si sceglie una trasformazione che porti il sistema dallo stato iniziale a uno stato con il volume finale attraverso un'adiabatica (in modo tale che in questo tratto la variazione di entropia sia nulla). Quindi si valuta la pressione raggiunta imponendo che pV^γ resti costante e, a questo punto, basta scrivere la variazione di entropia provocata dalla trasformazione isocora che porta il sistema nello stato finale dato.

10.6 Interpretazione microscopica dell'entropia

Come già abbiamo accennato all'inizio di questo capitolo, se riprendiamo con una telecamera certi processi elementari come l'urto tra due bocce o l'oscillazione di un pendolo, questi processi appaiono del tutto plausibili sia se osservati in avanti che all'indietro: questo riflette la perfetta simmetria delle leggi fisiche che li governano che sono identiche se si cambia il segno della variabile tempo. Si dice che le leggi fisiche sono **invarianti per inversioni temporali**.

Altri processi, però, non hanno la stessa caratteristica: se osservassimo all'indietro il filmato dello scioglimento di un cubetto di ghiaccio, del raffreddamento di un corpo caldo posto a contatto con un corpo freddo o della fuoriuscita di anidride carbonica da una lattina appena stappata, ci accorgeremmo immediatamente del fatto che quei processi non appaiono affatto plausibili!

La differenza rispetto ai processi elementari sopra illustrati sta nel fatto che nel primo caso i protagonisti dei filmati si possono considerare come particelle puntiformi, almeno entro certi limiti, mentre nel caso dei secondi non è così: il processo di scioglimento

del ghiaccio lo possiamo interpretare assumendo che il cubetto sia formato da tante goccioline d'acqua (puntiformi) che stanno vicine le une alle altre che, quando raggiungono la temperatura di zero gradi centigradi, cominciano a sentirsi meno legate alle altre e perciò sono più libere di muoversi. Analogamente possiamo pensare a due corpi a temperatura differente come composti da moltissime particelle più piccole che si agitano in modo tale che la loro energia cinetica media aumenti all'aumentare della temperatura del corpo: in questo modo il corpo caldo sarà composto di particelle che si muovono più rapidamente, in media, delle particelle del corpo freddo. Quando una particella del corpo caldo urta una del corpo freddo potrebbe trasferirle un po' di energia cinetica e provocarne così l'aumento di temperatura. Infine, nel caso dell'espansione libera del gas contenuto nella lattina appena stappata, evidentemente possiamo pensare che il gas sia formato da un gran numero di particelle che si muovono costrette a rimanere in un volume relativamente piccolo: nel momento in cui la lattina si stappa le particelle che si stavano muovendo in direzione del tappo non vengono più fermate da questo e continuano il loro moto verso l'esterno.

Questo però non spiega perché i processi avvengono solamente in uno dei possibili versi del tempo: in fondo i processi elementari che stanno alla base della descrizione che ne abbiamo dato sono perfettamente **reversibili**. Possono avvenire, cioè, in entrambi i sensi! Per esempio, nel caso del gas contenuto nella lattina, è vero che le particelle che stavano nella lattina e che si muovevano in direzione del tappo, una volta rimosso questo possono allontanarsi indisturbate, ma è anche vero che ci saranno altre particelle che dall'esterno della lattina si muovevano in direzione del tappo e che perciò, una volta rimosso quest'ultimo, dovrebbero penetrare nella lattina stessa. Forse succede, ma perché quelle che escono sono molte di più di quelle che entrano?

I processi **irreversibili**, che sono poi quelli che avvengono spontaneamente, hanno tutti la caratteristica, oltre che di procedere in un solo verso nel tempo, di provocare un aumento dell'entropia dell'Universo. Che ci sia una correlazione? Forse è il



Figura 10.2 Una particella si può disporre nella parte sinistra o nella parte destra di un volume.

fatto che l'entropia dell'Universo non può mai diminuire a determinare il fatto che il tempo scorre sempre in un'unica direzione? E se sí, perché?

Nel paragrafo precedente abbiamo calcolato la variazione di entropia conseguente al verificarsi di una trasformazione isoterma, eq. (10.56) di un gas, che riportiamo per comodità:

$$\Delta S_T = nR \log \frac{V_f}{V_i}. \quad (10.68)$$

che si può anche riscrivere come

$$\Delta S_T = NK \log \frac{V_f}{V_i}. \quad (10.69)$$

dove N è il numero di particelle di cui è composto il gas che subisce la trasformazione e k la **costante di Boltzmann**. Ora, proviamo ad analizzare il comportamento del gas dal punto di vista microscopico. Un gas perfetto è, per definizione, un gas composto di particelle che non interagiscono in alcun modo e prive di volume (puntiformi) quindi possiamo rappresentarlo come N particelle che si possono disporre in un numero arbitrario in qualsiasi volume. In particolare disponiamo le particelle in un volume V diviso idealmente in due parti V_1 , V_2 , con $V = V_1 + V_2$.

Cominciamo a considerare il caso $N = 1$, cioè di un gas (un po' speciale) costituito di un'unica particella. In questo caso ci sono solo due possibilità: la particella si trova o nel volumetto di sinistra oppure in quello a destra (Fig. 10.2).

Quando $N = 2$ le due particelle si possono disporre in più modi: possono stare entrambe a sinistra o

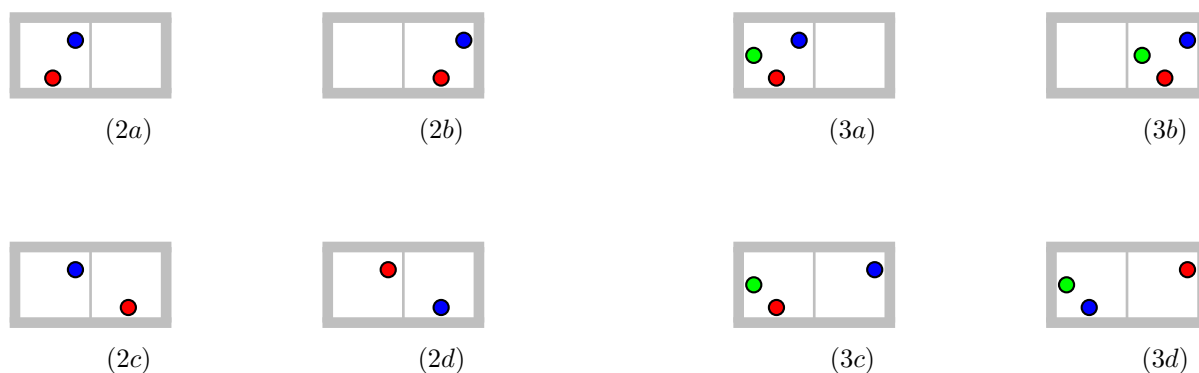


Figura 10.3 Ci sono quattro modi diversi di disporre 2 particelle nel volume.

entrambe a destra, ma si possono anche disporre in modo che la prima (quella blu in Fig. 10.3) sia a sinistra e l'altra (rossa) a destra o viceversa.

S poi $N = 3$ ci sono ben otto modi in cui possiamo disporre le particelle nelle due porzioni in cui è diviso il volume, come si vede dalla Fig. 10.4.

Molte di queste configurazioni (che chiameremo **microstati**), però, sono del tutto equivalenti l'una all'altra perché le particelle che compongono il gas sono tutte indistinguibili, quindi esistono configurazioni diverse che però producono gli stessi effetti macroscopici. Ad esempio, quando si verifica il microstato (2a) di Fig. 10.3 possiamo pensare, se la pressione esercitata sulle pareti della scatola è proporzionale al numero di particelle in vicinanza di ciascuna parete, che la pressione sul lato sinistro della scatola, esercitata dalle particelle presenti, sia maggiore rispetto a quella esercitata sul lato destro; nel caso (2b) sarà maggiore la pressione sul lato destro rispetto al lato sinistro, mentre nei casi (2c) e (2d) la pressione sarà la stessa in entrambi i casi. Se coloriamo le palline che rappresentano le singole particelle nello stesso modo si vede subito quali sono i microstati tra loro equivalenti, che producono, cioè, gli stessi **macrostati**. La Figura 10.5 mostra le stesse configurazioni di Fig. 10.3, ma impedendo di distinguere le singole particelle l'una dall'altra.

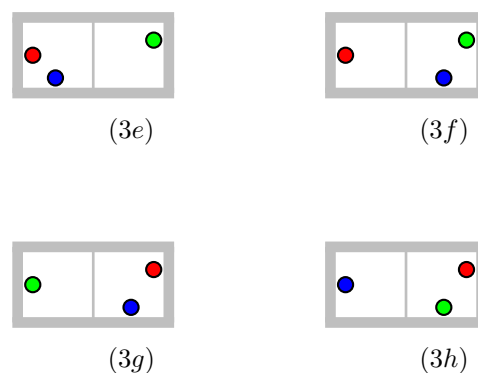


Figura 10.4 Se le particelle sono tre esistono otto diverse disposizioni.

In questo caso sono tre le configurazioni che producono effetti macroscopici diversi, ma una si può verificare in due modi diversi, quindi ha il doppio di probabilità di verificarsi rispetto alle altre due. Una cosa simile accade quando le particelle sono tre: il microstato (3a) e quello (3b) danno origine a due macrostati diversi (pressione maggiore a sinistra e a destra, rispettivamente); i macrostati da (3c) a (3e) sono tutti equivalenti (la pressione a sinistra è $2/3$ di quella a destra) così come quelli da (3f) a (3h). Ci sono quindi quattro macrostati diversi.

Proseguendo nell'esercizio si vede che il numero di modi in cui possiamo disporre quattro particelle è 16, con $N = 5$ avremmo 32 possibili combinazioni e così via. In effetti il numero di microstati possibili

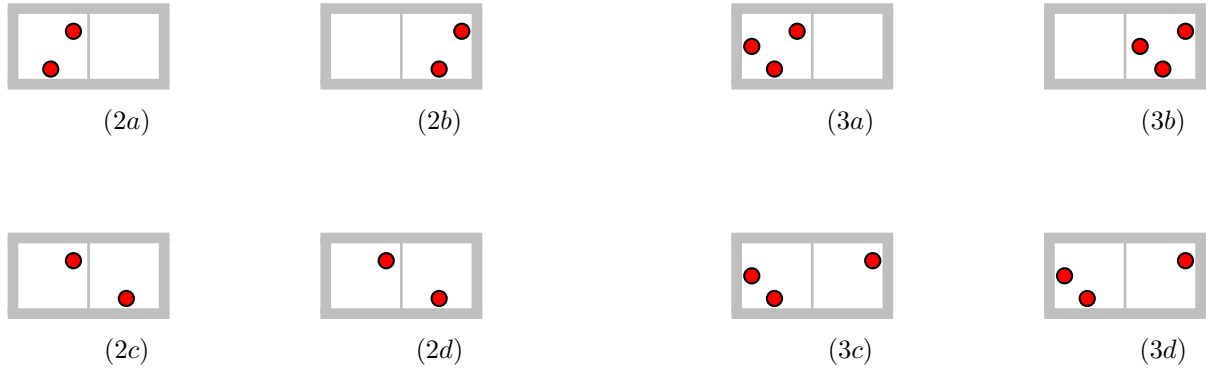


Figura 10.5 Ci sono quattro modi diversi di disporre 2 particelle nel volume, ma se le particelle sono indistinguibili le configurazioni davvero diverse sono solo tre.

è $\mu = 2^N$ perché ogni particella si può disporre in soli due modi: a destra o a sinistra; quindi N si possono disporre in $\mu = 2^N$ modi diversi. Se avessimo diviso il volume in tre parti, il numero di possibili configurazioni sarebbe stato pari a $\mu = 3^N$. È facile generalizzare dicendo che, se dividiamo il volume V in volumetti v possiamo scrivere che il numero di volumi a disposizione è V/v e quindi

$$\mu = \left(\frac{V}{v} \right)^N \quad (10.70)$$

Il logaritmo di quest'espressione, moltiplicato per la costante k dà

$$k \log \mu = Nk \log \frac{V}{v} \quad (10.71)$$

che somiglia moltissimo al secondo membro dell'equazione (10.56). La somiglianza è resa ancor più evidente se consideriamo una trasformazione del gas rappresentato dalle N particelle nella quale il volume passa da un valore V_i a un valore V_f . Il logaritmo del numero di possibili modi di disporre le N particelle, moltiplicato per k , vale

$$\mu_i = Nk \log \frac{V_i}{v} \quad (10.72)$$

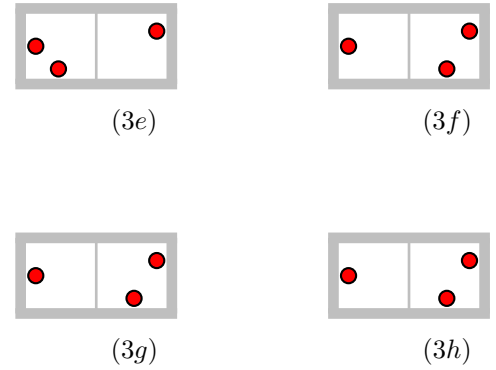


Figura 10.6 Il numero di macrostati con $N = 3$ è pari a quattro.

mentre nello stato finale vale

$$\mu_f = Nk \log \frac{V_f}{v} \quad (10.73)$$

e la variazione di μ , $\Delta\mu$ si scrive

$$\Delta\mu = \mu_f - \mu_i = Nk \left(\log \frac{V_f}{v} - \log \frac{V_i}{v} \right) \quad (10.74)$$

che sfruttando le proprietà dei logaritmi si riscrive come

$$\Delta\mu = Nk \log \frac{V_f}{V_i} = Nk \log \frac{V_f}{V_i}. \quad (10.75)$$

A questo punto è evidente che $\Delta\mu = \Delta S$ e quindi possiamo affermare che la variazione di entropia

di un sistema è una misura (un po' strana, ma pur sempre una misura) di quanto varia il numero di possibili modi in cui si possono disporre le particelle di cui è composto. Se l'entropia aumenta significa che il numero di modi in cui si possono disporre le particelle di cui è composto il sistema che stiamo esaminando aumenta. In un gas con un certo volume V ciascuna delle N particelle di cui è composto può occupare uno qualunque dei volumetti $v = V/N$ in cui possiamo dividere il volume totale (se il gas fosse ideale tutte le particelle potrebbero anche stare nello stesso volumetto). All'aumentare del volume aumenta anche il numero di volumetti a disposizione e quindi aumentano le possibili *scelte* che le particelle possono fare per andare a disporsi dentro il volume V . Quando invece un corpo si scalda, le particelle di cui è composto cominciano a muoversi più rapidamente. Questo moto non è un moto ordinato: alcune particelle avranno un'energia cinetica maggiore rispetto a quella media, altre l'avranno minore. Più è grande l'energia cinetica media (e quindi la temperatura), maggiore è l'ampiezza della distribuzione di energie cinetiche, quindi all'aumentare della temperatura aumentano le possibilità di ciascuna particella di *occupare* uno stato con una determinata velocità e di conseguenza aumenta l'entropia. In sostanza l'entropia misura il grado di **disordine** di un sistema: più il sistema è disordinato più è alta l'entropia di quel sistema, indipendentemente dalla natura del disordine. Per questa ragione l'entropia è una funzione di stato: che il sistema subisca un aumento di volume o un aumento di temperatura, il risultato è lo stesso e cioè che il numero di stati (posizione e velocità) che le particelle possono occupare aumenta.

La natura statistica dell'entropia spiega anche perché questa aumenta sempre nelle trasformazioni spontanee: il fatto è che, in linea di principio, le trasformazioni spontanee possono avvenire sia un senso (quello dell'aumento dell'entropia) che nell'altro (quello della sua diminuzione), ma il numero di stati che danno luogo a un aumento è sempre molto maggiore rispetto a quello degli stati che portano a una diminuzione dell'entropia. Il motivo per cui i processi procedono in un solo verso nel tempo è

dunque puramente statistico: l'aria in una stanza ne occupa tutto il volume semplicemente perché è improbabile che ne occupi solo la metà.

Lo si può vedere subito considerando i **macro-stati** del nostro gas di 2 o 3 particelle, illustrati nelle Figure 10.5 e 10.6. Nel primo caso gli stati in cui il gas occupa il volume uniformemente erano due su quattro, ma già nel secondo gli stati in cui la distribuzione delle particelle è più uniforme sono sei su otto. Non è difficile immaginare che all'aumentare delle particelle la probabilità di raggiungere uno stato macroscopico più uniforme aumenta vertiginosamente.

10.7 La media e la varianza di una distribuzione

Uno degli errori più frequenti in statistica consiste nel ritenere che se una variabile casuale ha una certa distribuzione, la somma di queste variabili avrà la stessa distribuzione. Per esempio, se lanciamo un dado abbiamo la stessa probabilità di ottenere un punteggio compreso tra 1 e 6 e possiamo affermare che la probabilità di ottenere il punteggio i , con $i = 1, \dots, 6$ è $p_i = \frac{1}{6}$. Se chiamiamo x_i il risultato ottenuto con il lancio i -esimo, contiamo il numero di volte $N(i)$ che si ottiene il punteggio i e ne facciamo un grafico in funzione dei possibili valori di x_i (che poi vanno da 1 a 6) otteniamo quella che si chiama una **distribuzione uniforme**: una distribuzione per cui la probabilità che esca il valore i è costante e pari a $1/n$, quando n è il numero di possibili valori che può assumere i (nel caso del dado $n = 6$). La Figura 10.7 è stata prodotta lanciando $N = 1\,000$ volte un dado (non serve farlo realmente, basta usare la capacità di un computer di generare numeri a caso). Come si vede la **frequenza** con la quale si presenta ciascun risultato è abbastanza costante (le fluttuazioni che si osservano sono di natura statistica e diventano via via più piccole al crescere di N).

Se ora, invece di lanciarne uno, di dadi ne lanciamo due possiamo ottenere tutti i valori compresi tra 2 e 12 e talvolta si è tentati di affermare che anche la

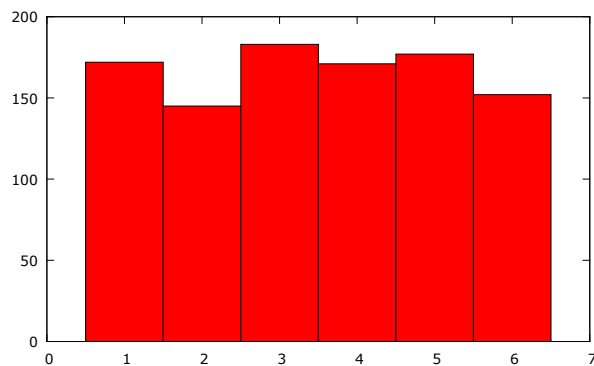


Figura 10.7 I lanci di un dado seguono una distribuzione uniforme.

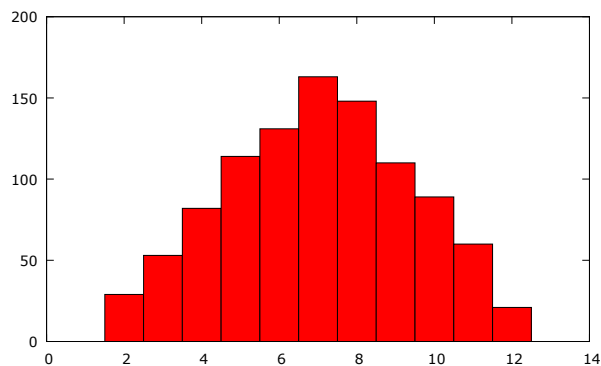


Figura 10.8 I lanci di due dadi seguono una distribuzione triangolare.

distribuzione di questi punteggi sarà uniforme, ma non è così. Il valore 2 (o il valore 12) si possono ottenere con un'unica combinazione: 1+1 (6+6 nel caso di 12). La combinazione 1+1 si può ottenere solo se entrambi i dadi danno come risultato 1. Il valore 7 invece si può ottenere con le combinazioni 1+6, 2+5, 3+4, ciascuna delle quali si può ottenere in due modi distinti (per esempio 1+6 = 6+1). Quindi il valore 7 si presenterà con una probabilità sei volte maggiore rispetto a 2 o a 12⁵. Il risultato è una distribuzione triangolare come quella di Fig. 10.8.

Se il numero di dadi di cui si somma il punteggio aumenta fino a diventare dieci, la distribuzione assume un aspetto ancor più piccato, come si vede nella Figura 10.9

Man mano che il numero di variabili casuali che si sommano cresce la distribuzione tende sempre di più a somigliare a una funzione **gaussiana**, dalla caratteristica forma a campana. È il **teorema del limite centrale** secondo il quale la somma di variabili casuali x , all'aumentare degli addendi, tende ad assumere una distribuzione gaussiana, qualunque

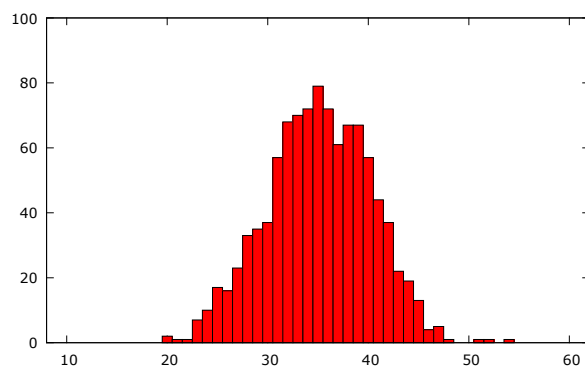


Figura 10.9 I lanci di dieci dadi seguono una distribuzione a campana centrata sul valore 35.

sia la distribuzione della variabile x . Ad esempio, per $N = 1000$ i punteggi potenziali vanno da 1000 a 6000, ma guardando la Figura ?? si capisce che la maggior parte dei valori non escono praticamente mai! Solo un piccolo gruppo di valori compresi attorno al valor medio di 3500 ha una certa probabilità di uscire.

L'intervallo di valori che si presentano con una frequenza ragionevolmente alta si può caratterizzare con la **deviazione standard** della distribuzione. Se il valor medio $\langle x \rangle$ si ottiene con la formula

$$\langle x \rangle = \frac{1}{N} \sum_{i=1}^N x_i, \quad (10.76)$$

⁵Il calcolo è semplice: per ottenere la combinazione 1+1 è necessario che il punteggio del primo dado sia 1, il che avviene con una probabilità di $1/6$, e che quello del secondo dado sia 1; questo avviene con probabilità $1/6$ perciò la probabilità complessiva è $\frac{1}{6} \times \frac{1}{6} = \frac{1}{36}$. Il punteggio 7 invece si ottiene con uno qualunque dei punteggi da 1 a 6, che si ottiene con probabilità 1 , seguito da uno solo dei possibili 6 risultati, che avviene con probabilità $1/6$. Di conseguenza la probabilità totale di ottenere 7 è $1 \times \frac{1}{6} = \frac{1}{6}$.

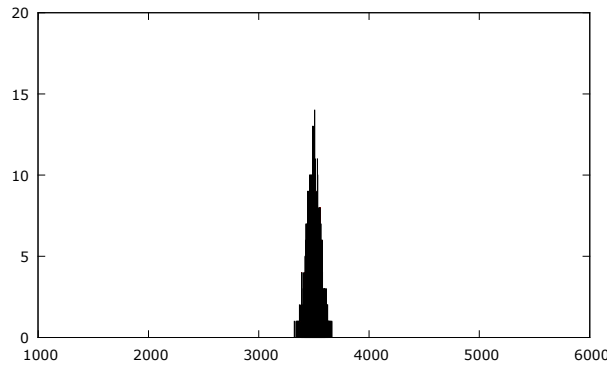


Figura 10.10 I lanci di mille dadi seguono una distribuzione praticamente gaussiana.

la deviazione standard σ deve fornire una misura di quanto distano i dati con probabilità elevata da questo numero. Una possibile misura di quanto il dato i -esimo disti da $\langle x \rangle$ sarebbe

$$d_i = |x_i - \langle x \rangle|. \quad (10.77)$$

L'operazione di **modulo** però è scomoda da manipolare e conviene usare la radice del quadrato della differenza, che fornisce lo stesso risultato

$$d_i = \sqrt{(x_i - \langle x \rangle)^2}. \quad (10.78)$$

Una misura complessiva, che dia il senso di quanto sia ampia la distribuzione, si può ottenere facendo la media delle *distanze* d_i :

$$\sigma = \frac{1}{N} \sqrt{(x_i - \langle x \rangle)^2}. \quad (10.79)$$

Nei lanci che abbiamo effettuato per produrre le figure riportate in questo paragrafo abbiamo ottenuto i valori di σ riportati in Tabella 10.1.

Il valore 1.7 per $N = 1$ indica che i dati si distribuiscono attorno al valor medio in modo da aggregarsi a una distanza (simmetrica a sinistra e a destra del valor medio) pari a 1.7, quindi attorno ai punti $3.5 - 1.7 = 1.8$ e $3.5 + 1.7 = 5.2$ che grosso modo stanno in mezzo alla parte sinistra e destra della distribuzione. Quando $N = 1000$, i dati si distribuiscono in modo da avere una distanza dal valor

N	σ	$\sigma/(b-a)$
1	1.7	0.28
2	2.4	0.20
10	5.2	0.09
1 000	17	0.03

Tavola 10.1 I valori della deviazione standard delle distribuzioni di 1000 lanci di N dadi. Nell'ultima colonna sono riportati i valori di σ divisi per l'ampiezza dei possibili valori $(b-a)$.

medio di 3 500 pari appena a 17. Se si misurano queste distanze in unità d'ampiezza dell'intervallo $b-a$ dei possibili valori si ottengono valori sempre più piccoli, il che indica che la distribuzione si stringe sempre più attorno al valor medio.

La statistica insegna che il valore di $\sigma/(b-a)$ diminuisce come $1/\sqrt{N}$, per cui nel caso di $N = 1000$ otteniamo $\sigma/(b-a) \simeq 1/\sqrt{1000} \simeq 0.03$. Quando il numero di dadi diventa comparabile con il **Numero di Avogadro** $N \simeq 6 \times 10^{23}$ la deviazione standard diventa dell'ordine di

$$\frac{\sigma}{b-a} \simeq \frac{1}{\sqrt{6 \times 10^{23}}} \simeq 1.3 \times 10^{-12}. \quad (10.80)$$

Questo significa che a fronte di un numero enorme di possibili configurazioni solo una su mille miliardi si verifica effettivamente con una certa probabilità. Questo spiega perché i gas si comportano in modo da occupare lo stato più uniforme possibile che è quello medio in cui la densità di gas è la stessa in ogni porzione di volume: le altre configurazioni sono possibili, ma estremamente improbabili. Di fatto, impossibili da osservare.

La teoria della Relatività Ristretta

PREREQUISITI: cinematica, dinamica elementare

Dopo gli esperimenti di **Michelson** e **Morley** non c'erano più scuse: per quanto possa sembrare strano la velocità della luce risulta essere indipendente dal moto relativo tra la sorgente e l'osservatore. Quest'osservazione sperimentale pone un serio problema filosofico. Secondo la *logica* comune, andando incontro a un raggio di luce ad altissima velocità, dovremmo vederlo arrivare verso di noi a una velocità pari alla somma tra la velocità del raggio di luce e la nostra; viceversa, se inseguissimo un raggio di luce alla sua stessa velocità dovremmo vederlo *fermo*! E se potessimo viaggiare a velocità più alte dovremmo riuscire addirittura a superarlo e a vederlo allontanarsi da noi in direzione opposta rispetto a quella nella quale si sta effettivamente muovendo rispetto alla sorgente considerata ferma! Eppure, secondo i risultati degli esperimenti non è così: se accendiamo una lampada vediamo la luce allontanarsi dalla lampada a una velocità di circa 300 000 km/s; se inseguissimo questa luce a bordo di un mezzo che si muove a velocità altissima, diciamo 299 000 km/s, non la vedremmo allontanarsi da noi a 1000 km/s, ma sempre a 300 000 km/s. Sembra una cosa completamente assurda! Impossibile! Un controsenso!

In un testo teatrale di **Luigi Malerba** intitolato "Stazione zero"¹ qualcuno dice che "Con i controsensi qualche volta si risolvono problemi che non si possono risolvere con la logica".

Nel 1905 **Albert Einstein** risolse il problema apparentemente irrisolvibile logicamente con quello che

¹Tratto da "Ai poeti non si spara", di Luigi Malerba, ed. Manni.

appariva essere un vero e proprio controsenso, che tuttavia è tale solo a causa della nostra esperienza. L'Universo ha le sue Leggi, che naturalmente potrebbero essere di qualunque natura e non c'è nessun motivo di principio secondo il quale l'Universo debba funzionare come a noi sembrerebbe logico! Se a noi le Leggi dell'Universo appaiono illogiche non è un problema dell'Universo: è un problema nostro! La fisica è una scienza sperimentale: ha a che fare con ciò che si misura, che ci piaccia o no. Una volta stabilito **sperimentalmente** che la velocità della luce è indipendente dal sistema di riferimento nel quale si esegue la misura, non possiamo far altro che accettarne le conseguenze, anche se possono sembrare dei *controsensi*.

La teoria della relatività ristretta (o **relatività speciale**) riguarda i sistemi di riferimento in moto l'uno rispetto all'altro di moto rettilineo uniforme.

11.1 Le trasformazioni di Lorentz

Consideriamo una sorgente di luce puntiforme nell'origine di un sistema di riferimento in quiete. Se la luce si muove a velocità c , preso un punto di coordinate (x, y, z) , il tempo necessario alla luce per arrivarci sarà

$$t = \sqrt{\frac{x^2 + y^2 + z^2}{c^2}}. \quad (11.1)$$

Prendendo il quadrato di quest'equazione e moltiplicando tutto per c^2 si ottiene

$$x^2 + y^2 + z^2 = c^2 t^2. \quad (11.2)$$

Se osservassimo lo stesso fenomeno da un sistema di riferimento che si muove lungo l'asse x a velocità v rispetto al primo, secondo la relatività galileiana l'equazione si scriverebbe

$$x'^2 + y'^2 + z'^2 = c^2 t'^2, \quad (11.3)$$

dove

$$\begin{cases} x' = x - vt \\ y' = y \\ z' = z \\ t' = t \end{cases} \quad (11.4)$$

e avremmo che

$$(x - vt)^2 + y^2 + z^2 = c^2 t^2 \quad (11.5)$$

cioè che

$$x^2 + v^2 t^2 - 2xvt + y^2 + z^2 = c^2 t^2 \quad (11.6)$$

in aperto contrasto con quanto scritto nell'equazione (11.2). È ovvio perché: nella relatività galileiana è impossibile che la velocità di qualcosa sia la stessa se misurata in due sistemi di riferimento in moto l'uno rispetto all'altro. Se vogliamo che gli osservatori in un sistema di riferimento e nell'altro siano in accordo circa le misure che conducono è necessario che x , y , z e t si trasformino in maniera diversa passando dall'uno all'altro sistema. Osservando le due equazioni è chiaro che, nel caso in esame, né y né z subiscono alcuna trasformazione e quindi l'unica maniera di far tornare le equazioni consiste nel trasformare t e x in modo tale da cancellare i termini indesiderati nell'equazione. La trasformazione deve essere tale per cui, per velocità non confrontabili con quella della luce ($v \ll c$), la trasformazione di x deve tornare a essere quella galileiana perciò possiamo provare² a mantenerla ponendo

$$x' = x - vt \quad (11.7)$$

²Se fossimo abili matematici scriveremmo subito che la trasformazione giusta deve essere lineare perciò avremmo che $x' = \alpha(x - \beta t)$ e $t' = \gamma(t - \delta x)$, determinando α , β , γ e δ imponendo l'invarianza della velocità della luce.

e di conseguenza t si deve trasformare in modo simile e cioè come

$$t' = t - \alpha x. \quad (11.8)$$

La trasformazione che abbiamo appena ipotizzato non potrà essere giusta, ma è un primo tentativo per trovare il modo di scriverla correttamente. Sostituiamo nell'equazione (11.3), ricordando che $y' = y$ e $z' = z$, per ottenere

$$(x - vt)^2 + y^2 + z^2 = c^2 (t - \alpha x)^2. \quad (11.9)$$

Espandiamo i quadrati:

$$(x^2 + v^2 t^2 - 2xvt) + y^2 + z^2 = c^2 (t^2 + \alpha^2 x^2 - 2\alpha xt) \quad (11.10)$$

e raccogliamo i termini simili:

$$x^2 (1 - c^2 \alpha^2) + y^2 + z^2 = t^2 (c^2 - v^2) - 2xt (c^2 \alpha - v). \quad (11.11)$$

Se vogliamo che quest'equazione sia uguale a (11.2), sicuramente non deve esserci il termine proporzionale a xt e dobbiamo evidentemente imporre che

$$c^2 \alpha = v \quad (11.12)$$

e quindi che $\alpha = v/c^2$. Così facendo l'equazione diventa

$$x^2 \left(1 - \frac{v^2}{c^2}\right) + y^2 + z^2 = c^2 t^2 \left(1 - \frac{v^2}{c^2}\right). \quad (11.13)$$

A questo punto, perchè il gioco sia fatto, è sufficiente che sia x che t si trasformino, passando da un sistema all'altro, in modo tale che i termini all'interno delle parentesi spariscano dall'ultima equazione scritta. Per questo dobbiamo fare in modo che le trasformazioni siano tali da contenere un fattore che cancella la parentesi $(1 - v^2/c^2)$, cioè dev'essere

$$x' = \frac{1}{\sqrt{1 - \frac{v^2}{c^2}}} (x - vt). \quad (11.14)$$

In questo modo, elevando al quadrato la coordinata x' compare un fattore a denominatore uguale a quello che compare nell'equazione (11.11), quindi questi due termini si semplificano e l'equazione assume la forma voluta. Una cosa analoga accade per t' . In definitiva le trasformazioni giuste sono

$$\begin{cases} x \rightarrow x' = \gamma(x - \beta ct) \\ t \rightarrow t' = \gamma\left(t - \frac{\beta}{c}x\right) \end{cases} \quad (11.15)$$

dove γ e β sono, rispettivamente

$$\gamma = \frac{1}{\sqrt{1 - \beta^2}} \quad \text{e} \quad \beta = \frac{v}{c}. \quad (11.16)$$

Osserviamo che per $v \ll c$, $\beta \simeq 0$ e $\gamma \simeq 1$ e le trasformazioni (11.15) si riducono a quelle di Galileo, perché v^2/c^2 risulta trascurabile rispetto a 1 e quindi $\gamma \simeq 1$. Inoltre $v/c^2 \simeq 0$ e quindi $t' \simeq t$.

Le trasformazioni (11.15) si chiamano **trasformazioni di Lorentz**, dal nome del fisico **Hendrik Lorentz** che aveva scoperto che queste stesse trasformazioni lasciavano invariate le **equazioni di Maxwell** passando da un sistema di riferimento inerziale assoluto (al tempo chiamato **etere**) a un qualunque sistema di riferimento in moto rettilineo uniforme rispetto a questo.

Le trasformazioni di Lorentz assumono una forma particolarmente simmetrica se si usano le **unità naturali**, cioè un sistema di unità di misura nel quale $c = 1$ ed è adimensionale. In questo sistema le velocità sono grandezze fisiche adimensionali e si misurano in frazioni della velocità della luce, mentre le lunghezze si misurano in velocità per tempo (ed essendo c adimensionale si misurano perciò in secondi)³. In questo sistema, infatti $\beta = v$ e

$$\begin{cases} x \rightarrow x' = \gamma(x - \beta t) \\ t \rightarrow t' = \gamma(t - \beta x) \end{cases} \quad (11.17)$$

Usare le unità naturali semplifica enormemente i conti, perché si eliminano tutti i fattori c a nume-

ratore e a denominatore, che rendono complicate le formule. Una volta trovato il risultato, per ottenere i valori delle misure in unità SI basterà moltiplicarlo per un'opportuna potenza di c , tale da renderlo dimensionalmente corretto. Ad esempio, nel caso delle trasformazioni (11.17), per tornare a quelle espresse nel SI, basta osservare che l'equazione che dà x' deve avere le dimensioni di una lunghezza. Essendo γ adimensionale, l'espressione tra parentesi deve essere una lunghezza. x lo è, mentre βt ha le dimensioni di un tempo. L'unico modo di far avere a questo addendo le dimensioni giuste consiste nel moltiplicarlo per una velocità che è chiaramente c , per ottenere $x' = \gamma(x - \beta ct)$. Allo stesso modo, l'espressione che ci dà t' deve avere le dimensioni di un tempo. Tra parentesi c'è la differenza tra t , che ha le giuste dimensioni, e βx , che ha le dimensioni di una lunghezza. Per far avere a questo addendo le dimensioni di un tempo occorre dividerlo per una velocità e quindi l'espressione diventa $t' = \gamma\left(t - \frac{\beta}{c}x\right)$. Nel resto di questo capitolo, ove non diversamente indicato, si usano le unità naturali per la derivazione delle relazioni tra le grandezze fisiche.

A differenza di Lorentz, Einstein aveva capito che non era necessario supporre l'esistenza di un etere con strane proprietà di trasformazione per spiegare il fatto sperimentale secondo il quale la luce si muove sempre alla stessa velocità, in qualunque sistema di riferimento. Secondo la visione di Einstein lo spazio e il tempo non sono concetti assoluti, come fino ad allora si era ritenuto, ma essendo essi stessi grandezze fisiche misurabili, erano concetti relativi: per eseguire una misura bisogna sempre confrontare una grandezza fisica con una ad essa omogenea. Spazio e tempo non fanno eccezione: se gli strumenti attraverso i quali li misuro cambiano passando da un sistema all'altro, la misura dello spazio e del tempo non è assoluta.

11.2 La dilatazione del tempo

Ma che vuol dire che il tempo non è assoluto? Il tempo, come abbiamo detto, è una grandezza fisica e pertanto occorre misurarla con un qualche stru-

³L'anno luce, comunemente usato in astrofisica, è una misura di lunghezza in queste unità: rappresenta la lunghezza del percorso fatto dalla luce in un anno.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 11.1 Usando un orologio a luce si capisce bene che il tempo non può essere assoluto [https://www.youtube.com/watch?v=g2eI0Wi2bVA]

mento. Lo strumento con il quale si misura il tempo è l'orologio. Secondo la teoria di Einstein, detta della **relatività**, la durata di un secondo non è la stessa per tutti gli osservatori: quando per un osservatore fermo sono passati τ secondi, per un osservatore che si muove a velocità v rispetto a questo ne sono passati, usando lo stesso orologio

$$t = \gamma\tau = \frac{1}{\sqrt{1 - \beta^2}}\tau. \quad (11.18)$$

Nello scrivere questo tempo abbiamo semplicemente usato la trasformazione di Lorentz ponendo $x = 0$ (dal momento che possiamo scegliere sempre di mettere l'orologio nell'origine del sistema di riferimento dell'osservatore fermo).

Si vede subito che, per $v \ll c$ il fattore di Lorentz $\gamma \simeq 1$ e, come ci aspettiamo, il tempo appare quasi assoluto. Se però la velocità dell'osservatore diventa ragguardevole si possono verificare strani fenomeni (strani, naturalmente, per quel che è la nostra esperienza quotidiana).

Il tempo τ misurato con un orologio fermo nello stesso sistema di riferimento in cui si esegue la misura di un qualche fenomeno fisico si chiama **tempo proprio**.

Qui bisogna fare attenzione a non commettere errori grossolani: se due osservatori sono in moto relativo rettilineo e uniforme l'uno rispetto all'altro, per ciascuno di essi il tempo scorre esattamente come ci si aspetta che scorra. Se però un osservatore misura il tempo con un orologio che si trova a bordo dell'altro sistema vede una discrepanza rispetto a quel che misura con il proprio orologio! È evidente che, dal momento che il moto è relativo, l'osservatore in moto rispetto al primo vedrà l'orologio del primo

muoversi rispetto a lui a velocità v e ne concluderà che quell'orologio va indietro rispetto a quello che porta con sé. In maniera del tutto analoga, l'osservatore fermo, osservando l'orologio che porta con sé l'osservatore in moto, lo vedrà rallentare rispetto al suo. Entrambi gli osservatori trarranno le stesse conclusioni, senza ambiguità. Osserviamo anche che la distinzione tra osservatore fermo e in moto è del tutto arbitraria: non c'è modo di stabilire chi si stia muovendo e chi sta fermo!

Ci si potrebbe chiedere se l'effetto è reale o apparente, ma a dir la verità anche questa domanda appare mal formulata: la fisica è una scienza sperimentale ed è reale quel che si misura. Se attraverso le misure trovo che il tempo scorre in maniera diversa secondo il sistema di riferimento nel quale faccio la misura, ne devo concludere che è così. Punto!

Esercizio 11.1 *Il tempo in orbita*

I satelliti della costellazione GPS orbitano attorno alla Terra, a una quota di 20 000 km, muovendosi a una velocità media di circa 4 km/s. Calcola la durata di un secondo, di un giorno, di un mese e di un anno a bordo del satellite, se il tempo è misurato con un orologio fermo a Terra [10].

SOLUZIONE →

In ogni caso l'effetto è molto più *reale*⁴ di quanto si pensi. Il nostro pianeta è oggetto di una continua pioggia di particelle che provengono dallo spazio, i così detti **raggi cosmici**. Queste particelle, urtando gli atomi degli strati più alti dell'atmosfera, a una quota di circa 10 km, ne producono altre chiamate **muoni**. I muoni sono particelle instabili, che *decadono*, cioè si trasformano spontaneamente, in un elettrone e due **neutrini** mediamente in circa $2 \mu\text{s}$. Se i muoni si muovessero alla velocità della luce, in $2 \mu\text{s}$

⁴Riflettete sul significato di questa parola: per la fisica è reale ciò che si misura. Non ha senso chiedersi se effettivamente il tempo scorra più lentamente oppure se si tratti di una limitazione della nostra capacità di misurarlo. Quel che si misura è. E del resto il caso del decadimento del muone dimostra che la domanda è praticamente priva di senso.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 11.2 I navigatori GPS funzionano grazie alla teoria della relatività. In questo filmato se ne spiegano i principi di funzionamento. Risolvete l'esercizio 11.1 per valutare l'effetto della dilatazione dei tempi su questo sistema di navigazione [<https://www.youtube.com/watch?v=8hFhKDxfyHQ>].

percorrerebbero circa $2 \times 10^{-6} \times 3 \times 10^8 = 600$ m. In realtà si osservano numerosissimi muoni a livello del mare (un centinaio per metro quadro al secondo), quindi questi devono poter aver viaggiato per oltre 10 km. La teoria della relatività spiega questa apparente stranezza: nel sistema di riferimento del muone, nel quale è fermo, il tempo scorre in maniera tale che mediamente, trascorsi $2 \mu\text{s}$, buona parte dei muoni decadono. Ma quando noi osserviamo i muoni provenire dallo spazio, li vediamo muoversi a velocità molto vicine a quella della luce, per cui, per noi che siamo a terra, il tempo a bordo del muone scorre molto più lentamente. Occorrono $\gamma\tau$ secondi misurati a terra per far decadere i muoni e se γ è abbastanza grande questo tempo si può dilatare a dismisura. Questo fenomeno si osserva quotidianamente nei laboratori di fisica, agli acceleratori di particelle. E anche se non si vede direttamente, il fenomeno è ormai ben presente nella nostra vita quotidiana. I sistemi di navigazione GPS funzionano grazie a una costellazione di satelliti a bordo dei quali ci sono orologi atomici molto precisi. Se non si tenesse conto degli effetti della relatività nella marcia di questi orologi i sistemi di navigazione sarebbero del tutto inutili perché fornirebbero la posizione dell'utente con un errore crescente nel tempo, che può arrivare anche a diversi km!

11.3 Contrazione della lunghezza

Contestualmente al tempo, secondo la teoria della relatività, si modificano anche le proprietà dello spazio. In particolare, consideriamo una riga lunga ℓ in un sistema di riferimento fermo (ricordiamo però che si tratta di una convenzione, non potendo affatto stabilire se il sistema sia davvero fermo o si stia muovendo di moto rettilineo uniforme rispetto a un altro) orientata lungo l'asse x in modo tale che l'estremo sinistro sia nell'origine del sistema di riferimento e l'altro nel punto di coordinate $(\ell, 0, 0)$. Osservando questa stessa riga da un altro sistema di riferimento, in moto con velocità v rispetto al primo lungo l'asse x , la coordinata dell'estremo sinistro diventa

$$x'_L = -\gamma\beta t \quad (11.19)$$

mentre la coordinata dell'estremo destro diventa

$$x'_R = \gamma(\ell - \beta t). \quad (11.20)$$

La lunghezza della riga per l'osservatore in moto è dunque $x'_R - x'_L$ che vale

$$x'_R - x'_L = \ell' = \gamma(\ell - \beta t + \beta t) = \gamma\ell \quad (11.21)$$

L'osservatore in moto dunque leggerà il numero ℓ sulla riga a una distanza pari a

$$\ell = \frac{\ell'}{\gamma} \quad (11.22)$$

ed essendo sempre $\gamma \geq 1$ vedrà la riga contrarsi di un fattore γ rispetto al suo sistema di unità di misura.

Il lungo percorso fatto dai muoni visti nel Paragrafo 11.2 prima di decadere si può spiegare anche alla luce di questo fenomeno. Il muone in volo, infatti, nel suo sistema di riferimento, vede la Terra corrergli incontro a grandissima velocità. La distanza tra il muone e la Terra dunque appare al muone contratta di un fattore γ per cui nel suo sistema di

riferimento questa particella non percorre molto più dei 600 m previsti.

Dalle trasformazioni di Lorentz è anche evidente che le dimensioni trasversali degli oggetti (le lunghezze, cioè, misurate lungo direzioni ortogonali a quella nella quale si stanno muovendo l'uno rispetto all'altro i sistemi di riferimento) non cambiano quando si passa da un sistema all'altro.

Esercizio 11.2 *Il viaggio d'un muone*

In seguito alle interazioni dei [raggi cosmici](#) primari con i nuclei dei gas dell'atmosfera, a circa 10 km di quota, si produce un muone che viaggia verso terra a una velocità pari al 99.5 % di quella della luce. Quanto dovrà percorrere il muone prima di urtare la Terra, secondo il suo metro?

[SOLUZIONE](#) →

11.4 Composizione delle velocità

Dal momento che secondo la teoria della relatività ristretta le trasformazioni di Galileo non sono più valide, occorre trovare nuove trasformazioni per calcolare la velocità di un oggetto come vista da un sistema di riferimento in moto rispetto a un altro.

Farlo non è difficile: basta osservare che la velocità è data dal rapporto $u = \Delta x / \Delta t$ tra lo spazio percorso Δx e il tempo impiegato a percorrerlo Δt . Trasformando con Lorentz queste due quantità otteniamo, assumendo sempre che l'asse x sia orientato nella direzione del moto relativo tra i due sistemi, che

$$u' = \frac{\Delta x'}{\Delta t'} = \frac{\Delta x - \beta \Delta t}{\Delta t - \beta \Delta x}. \quad (11.23)$$

Dividendo l'ultimo membro per Δt si ottiene

$$u' = \frac{\Delta x'}{\Delta t'} = \frac{\frac{\Delta x}{\Delta t} - \beta}{1 - \beta \frac{\Delta x}{\Delta t}}. \quad (11.24)$$

Il paradosso dei gemelli

Uno dei più noti apparenti paradossi della fisica relativistica è quello detto *dei gemelli*. Ammettiamo che in un lontano futuro sia possibile costruire un'astronave superveloce a disposizione di due gemelli: Ulisse e Telemaco. Ulisse sale a bordo dell'astronave e intraprende un viaggio che, per lui, dura circa dieci anni. Al ritorno sulla Terra i due gemelli non sarebbero più tali perché, secondo la teoria della relatività, Telemaco da Terra vede scorrere il tempo del fratello molto più lentamente del suo. Quindi Ulisse invecchia più lentamente di Telemaco e torna a casa più giovane.

Il paradosso sta nel fatto che lo stesso discorso si potrebbe applicare a Ulisse, il quale nel suo sistema di riferimento vede il tempo scorrere come se fosse a Terra e quindi invecchia normalmente, ma dovrebbe vedere Telemaco sfrecciare a velocità elevatissime e quindi dovrebbe vedere il suo tempo rallentare. È quindi Telemaco a risultare più giovane, per Ulisse.

In realtà il paradosso così com'è non si può formulare perché la teoria della relatività ristretta si applica solo ai sistemi di riferimento inerziali, che si devono muovere di moto rettilineo uniforme l'uno rispetto all'altro. Se Ulisse a un certo punto torna indietro, il suo moto non può essere rettilineo uniforme (per non parlare delle fasi di arrivo e partenza), quindi in questo caso la relatività ristretta non vale. Vale però la [relatività generale](#) secondo la quale comunque avviene che Ulisse invecchia meno rapidamente di Telemaco.

e quindi

$$u' = \frac{u - \beta}{1 - \beta u}. \quad (11.25)$$

Osserviamo anzitutto che l'espressione risulta essere adimensionale, come deve, dal momento che in unità naturali le velocità non hanno dimensioni e si misurano in frazioni di velocità della luce. Dove si legge u , dunque, si deve sempre intendere *misurato*

in unità di c e quindi numericamente pari a u/c . In questo caso è utile riscrivere l'espressione in unità SI. Il passaggio a unità SI è quanto mai semplice: basta moltiplicare il tutto per una velocità, che non può che essere c . L'espressione in unità SI è dunque, ricordando che per u dobbiamo intendere u/c ,

$$u' = \frac{\frac{u}{c} - \beta}{1 - \beta \frac{u}{c}} c = \frac{u - v}{1 - \frac{uv}{c^2}}, \quad (11.26)$$

che è l'espressione che si trova comunemente scritta sui libri (ricordiamo che $\beta = v/c$). Quest'espressione, sebbene un po' più complessa, è sempre facile da ricordare, perché al numeratore c'è la differenza tra le velocità del punto materiale e del sistema di riferimento dal quale lo si guarda, come nella relatività galileiana, che deve essere corretta relativisticamente applicando un fattore che nel limite $v \ll c$ diventa pari a 1, che è quello al denominatore. Questo fattore deve essere adimensionale e deve dipendere da u e da v che, per natura, sono vettori. Un modo per costruire una grandezza scalare usando due vettori è farne il prodotto scalare $\mathbf{u} \cdot \mathbf{v} = uv$ nel caso in esame perché $\mathbf{u}$ e $\mathbf{v}$ sono tra loro paralleli. Questo prodotto ha le dimensioni di una velocità al quadrato e per renderlo adimensionale si può dividerlo per c^2 .

Per $v \ll c$ il rapporto uv/c^2 è piccolo e trascurabile rispetto a 1 e riotteniamo la trasformazione di Galileo. Per $v \simeq c$, $u' \simeq -c$. La velocità di qualunque cosa vista da un sistema di riferimento che si muove alla velocità della luce è sempre pari alla velocità della luce. Inoltre, se $u = c$, la velocità $u' = c$ per ogni valore di v (come ci aspettiamo dal momento che tutto deriva dalla solita osservazione sperimentale secondo la quale la velocità della luce è indipendente dal sistema di riferimento nel quale la si misura). Esiste dunque un **limite** alla velocità con la quale si possono muovere gli oggetti: nessuno potrà mai misurare una velocità superiore a quella della luce. Il risultato può apparire quanto meno sorprendente, ma per quanto strano possa sembrarci è un fatto sperimentale che, fin quando non verrà smentito con altre osservazioni, resta valido e dobbiamo metterci l'anima in pace. Del resto, che qualcosa si possa muovere a qualunque velocità,

in fondo, non è che un pregiudizio: nessuno ci è mai riuscito.

Un'altra osservazione utile è la seguente: supponiamo di non conoscere affatto le trasformazioni di Lorentz, ma di sapere che le velocità si trasformano in modo tale che, per v piccole, valgono le leggi di trasformazione di Galileo, mentre per v grandi le velocità devono tendere a quella della luce. Possiamo sempre scrivere che

$$u' = A(u - v) \quad (11.27)$$

dove A è un numero che dipende da u e da v e che deve tendere a 1 quando v è piccolo. Possiamo dunque scrivere $A = 1 + B$ con $B = f(u, v)$, tale che per v che tende a zero, $f(u, v) \simeq 0$. A questo punto possiamo ripetere il ragionamento fatto sopra: u e v sono per natura dei vettori, mentre B è uno scalare. Un modo di costruire uno scalare con due vettori è farne il prodotto scalare: $\mathbf{u} \cdot \mathbf{v} = uv$ nel caso specifico. Il prodotto uv non è adimensionale come deve essere B , quindi dobbiamo dividerlo per una velocità assoluta al quadrato: c^2 . Otteniamo

$$u' = \left(1 + \frac{uv}{c^2}\right)(u - v) \quad (11.28)$$

che è solo apparentemente diversa dalla relazione esatta ricavata sopra. Infatti, come si vede nell'Appendice al paragrafo sull'[approssimazione di funzioni](#), l'espressione $1/(1 - x)$ si può approssimare, per x piccolo, a $1 + x$ (vedi eq. (25.68)). Si vede subito che

$$u' = \frac{u - v}{1 - \frac{uv}{c^2}} \simeq \left(1 + \frac{uv}{c^2}\right)(u - v). \quad (11.29)$$

Come si vede si può ricavare un risultato abbastanza vicino a quello corretto semplicemente usando argomenti dimensionali e un po' di matematica.

11.5 I quadrivettori

La maniera in cui si descrivono lo spazio e il tempo suggerisce che la descrizione classica secondo la quale lo stato di un punto materiale è determinato

da posizione e velocità non è corretta, ma è un'approssimazione della descrizione corretta nel limite di basse velocità.

Lo stato di un oggetto non si può rappresentare con posizione e velocità, che sono due vettori nello spazio a tre dimensioni, perché le posizioni e le velocità non sono assolute, ma dipendono dall'osservatore. La grandezza fisica $s^2 = t^2 - (x^2 + y^2 + z^2)$ ⁵, al contrario, è assoluta: assume lo stesso valore in tutti i sistemi di riferimento. Questa grandezza quindi si chiama **invariante di Lorentz**.

Un punto materiale che si trova al tempo t alle coordinate (x, y, z) in un sistema di riferimento, in un altro sistema di riferimento che si muove rispetto al primo di moto rettilineo uniforme parallelamente all'asse x , avrebbe coordinate $(x', y', z') = (\gamma(x - \beta t), y, z)$ all'istante $t' = \gamma(t - \beta x)$. La quantità s^2 , espressa nel sistema di riferimento in moto, è

$$s'^2 = t'^2 - (x'^2 + y'^2 + z'^2) = \gamma^2 (t - \beta x)^2 - (\gamma^2 (x - \beta t)^2 + y^2 + z^2). \quad (11.30)$$

Svolgendo i quadrati l'espressione sopra scritta diventa

$$s'^2 = \gamma^2 (t^2 + \beta^2 x^2 - 2\beta xt) - \gamma^2 (x^2 + \beta^2 t^2 - 2\beta xt) - y^2 - z^2 \quad (11.31)$$

e, al solito, raccogliendo i termini comuni abbiamo

$$s'^2 = t^2 \gamma^2 (1 - \beta^2) - x^2 \gamma^2 (1 - \beta^2) - y^2 - z^2. \quad (11.32)$$

I termini misti proporzionali a xt si elidono a vicenda. Ora osserviamo che $\gamma^2 = (1 - \beta^2)^{-1}$ e i coefficienti di x^2 e di $2t^2$ valgono quindi entrambi 1. Perciò $s'^2 = s^2$. Si dice che s^2 è un **invariante relativistico** perché il suo valore non cambia passando da un sistema di riferimento a un altro.

Possiamo pensare a s^2 come al prodotto scalare di un vettore per sé stesso, se costruiamo dei vettori

⁵L'espressione di s^2 in unità SI è $s^2 = c^2 t^2 - (x^2 + y^2 + z^2)$

a quattro dimensioni con le posizioni e gli istanti di tempo in questo modo:

$$\mathbf{s} = (t, ix, iy, iz) \quad (11.33)$$

dove $i = \sqrt{-1}$ è l'unità immaginaria. In questo modo, prendendo il prodotto scalare $\mathbf{s} \cdot \mathbf{s}$, che si ottiene sommando i prodotti delle coordinate omologhe, otteniamo proprio il valore di s^2 .

Possiamo dunque pensare a $\mathbf{s}$ come a dei vettori in uno spazio quadridimensionale in cui le trasformazioni di Lorentz eseguono delle rotazioni⁶ di questi **quadrivettori**. La prima componente di questo quadrivettore è il tempo (moltiplicato per c se si scrive in unità SI), mentre le restanti coordinate sono quelle spaziali moltiplicate per l'unità immaginaria. Possiamo anche eliminare quest'ulteriore complicazione ridefinendo l'operazione di prodotto scalare in questo spazio particolare (detto **spazio di Minkowski**)⁷: basta ricordare che le componenti spaziali e quelle temporali vanno sommate col segno opposto (quale segno di fatto è irrilevante, tanto cambierebbe solo il segno dell'invariante, che tuttavia continuerebbe a restare costante).

In fisica relativistica, dunque, i concetti di **posizione** e di **velocità** di un punto perdono parte del loro significato, giacché la posizione e la velocità di un punto materiale sono qualcosa di relativo: dipendono dall'osservatore e dal tempo. Queste due nozioni sono sostituite dalla nozione di **evento**, che si caratterizza fornendo sia la posizione sia il tempo al quale la posizione del punto è stata raggiunta per un determinato osservatore.

11.6 Il quadrivettore energia-impulso

Possiamo definire altri quadrivettori a partire dal quadrivettore posizione che definisce un evento. Come nel caso dei comuni vettori, un quadrivettore moltiplicato scalarmente per un altro quadri-

⁶Le rotazioni di un vettore ne cambiano le coordinate, ma non il modulo.

⁷Dal nome del matematico **Hermann Minkowski**.

vettore dà uno scalare (che è quindi un invariante per trasformazioni di Lorentz), mentre un quadrivettore moltiplicato per uno scalare è ancora un quadrivettore.

In fisica classica si definisce la quantità di moto come il prodotto $\mathbf{p} = m_0 \mathbf{u}$ della massa m_0 di un punto materiale per la sua velocità $\mathbf{u}$. Per misurare una velocità dobbiamo prendere uno spostamento e dividerlo per un tempo. Consideriamo un punto materiale che al tempo t_1 del nostro orologio si trova in $\mathbf{r}_1 = (x_1, x_2, x_3)$. Per questo punto materiale possiamo definire il quadrivettore $\underline{r}_1 = (ct_1, x_1, y_1, z_1)$ (in questo caso stiamo usando le unità SI, perché questo ci aiuterà nel dare la corretta interpretazione ai risultati). Se il punto si trova in $\mathbf{r}_2 = (x_2, y_2, z_2)$ al tempo t_2 possiamo definire un secondo quadrivettore $\underline{r}_2 = (ct_2, x_2, y_2, z_2)$. La differenza tra due quadrivettori è ancora un quadrivettore:

$$\underline{\Delta r} = (c\Delta t, \Delta x, \Delta y, \Delta z) \quad (11.34)$$

dove $\Delta t = t_2 - t_1$ e $\Delta x = x_2 - x_1$ (e analogamente per le altre coordinate). Potremmo chiamare questo quadrivettore **spostamento**. Si sarebbe a questo punto tentati di dividere tutto per Δt per ottenere un quadrivettore che rappresenti la velocità, ma Δt non è uno scalare nello spazio di Minkowski, perché non è invariante per trasformazioni di Lorentz. Per trovare uno scalare che abbia le dimensioni di un tempo dobbiamo ricorrere al tempo proprio, che è ben definito ed è sempre lo stesso in ogni sistema di riferimento. Il tempo proprio in questo caso è il tempo misurato con un orologio che si muove insieme al punto materiale di cui si sta misurando la posizione. Dividendo il quadrivettore spostamento per lo scalare $\Delta\tau$ e ricordando che $t = \gamma\tau$, otteniamo un ulteriore quadrivettore che potremmo chiamare **quadrivelocità**:

$$\underline{v} = \frac{\underline{\Delta r}}{\Delta\tau} = \left(\frac{c\Delta t}{\Delta\tau}, \frac{\Delta x}{\Delta\tau}, \frac{\Delta y}{\Delta\tau}, \frac{\Delta z}{\Delta\tau} \right) = \gamma (c, u_x, u_y, u_z) \quad (11.35)$$

con $u_x = \frac{\Delta x}{\Delta t} = \frac{\Delta x}{\gamma\Delta\tau}$ la componente x della velocità tridimensionale (e analogamente per u_y e u_z). Osserviamo che $\gamma = (1 - \beta^2)^{-1/2}$ è il fattore di Lorentz

ottenuto prendendo come $\beta = u/c$ la velocità della particella in unità di c , perché stiamo scrivendo le grandezze trasformando il tempo misurato nel sistema di riferimento del laboratorio al sistema di riferimento in cui la particella di cui si sta misurando la velocità è ferma.

La quadrivelocità è un quadrivettore (essendo un quadrivettore diviso uno scalare) e dunque per essa valgono le **trasformazioni di Lorentz**. Il quadrivettore velocità $\underline{v}$, misurato in un sistema di riferimento fermo, visto da un sistema di riferimento in moto con velocità $\mathbf{V}$ rispetto al primo, si ottiene trasformando con Lorentz le componenti temporali e spaziali. Per semplicità supponiamo che, in unità naturali,

$$\underline{v} = \gamma (v_t, v_x, v_y, v_z) = \gamma (1, u, 0, 0) \quad (11.36)$$

e

$$\mathbf{V} = (\beta, 0, 0), \quad (11.37)$$

cioè che l'oggetto di quadrivelocità $\underline{v}$ e l'osservatore si muovano lungo lo stesso asse che consideriamo come l'asse x . Secondo le trasformazioni di Lorentz, dal sistema con velocità $\mathbf{V}$ vedremo il quadrivettore $\underline{v}$ avere le componenti

$$\begin{aligned} v'_t &= W (v_t - \beta v_x) = W\gamma (1 - \beta u) \\ v'_x &= W (v_x - \beta v_t) = W\gamma (u - \beta) \end{aligned} \quad (11.38)$$

con $W = 1/\sqrt{1 - \beta^2}$, che è il fattore di Lorentz del sistema in moto rispetto a quello fisso. Ora osserviamo che le componenti del quadrivettore $\underline{v}$ sono

$$\underline{v} = \left(\frac{\Delta t}{\Delta\tau}, \frac{\Delta x}{\Delta\tau}, 0, 0 \right) \quad (11.39)$$

e quelle di $\underline{v}'$

$$\underline{v}' = \left(\frac{\Delta t'}{\Delta\tau}, \frac{\Delta x'}{\Delta\tau}, 0, 0 \right). \quad (11.40)$$

Pertanto, la velocità misurata da un osservatore nel sistema in moto, è

$$\frac{\Delta x'}{\Delta t'} = \frac{\Delta x'}{\Delta \tau} \frac{\Delta \tau}{\Delta t'} = \frac{v'_x}{v'_t}. \quad (11.41)$$

Ritroviamo così la formula per la [composizione delle velocità](#)

$$\frac{\Delta x'}{\Delta t'} = u' = \frac{W\gamma(u - \beta)}{W\gamma(1 - \beta u)}. \quad (11.42)$$

Questo vi fa capire che l'uso della quadrivelocità può essere comodo quando si debbano ricavare le leggi di composizione della velocità perché è sufficiente ricordare che si tratta di un quadrivettore e le trasformazioni di Lorentz di questi, ma che bisogna stare attenti a interpretarne le componenti. La velocità di un corpo **non** è la componente spaziale di un quadrivettore velocità, ma il rapporto tra la sua componente spaziale e quella temporale!

Vale la pena osservare che la somma di due quadrivelocità, pur essendo un quadrivettore, non ha alcun significato fisico particolare e **non è una quadrivelocità**. Infatti, una proprietà della quadrivelocità abbastanza evidente è che il suo modulo vale c (oppure 1 in unità naturali). Infatti, facendo il quadrato del quadrivettore [11.35](#) si ottiene

$$\underline{v}^2 = \gamma^2 (c^2 - u^2) = \frac{1}{1 - \frac{u^2}{c^2}} (c^2 - u^2) = c^2. \quad (11.43)$$

Una quadrivelocità quindi è un quadrivettore il cui modulo vale sempre c e di conseguenza la somma di due quadrivelocità non può essere una quadrivelocità. In altre parole la quadrivelocità non è una grandezza fisica *interessante*: si tratta solo di un artificio per costruire altri quadrivettori d'interesse. In effetti, moltiplicando tutto per la massa m_0 del punto materiale che si è spostato nel tempo Δt da $\mathbf{r}_1$ a $\mathbf{r}_2$ costruiamo il quadrivettore

$$\underline{p} = m_0 \underline{v} = \gamma (m_0 c, m_0 u_x, m_0 u_y, m_0 u_z) = \gamma m_0 (c, \mathbf{u}). \quad (11.44)$$

Per come è stata costruita, la componente spaziale di questo quadrivettore dovrebbe rappresentare la

quantità di moto⁸ della particella. Vediamo invece che risulta essere pari a $\gamma m_0 \mathbf{u}$. C'è un fattore γ di troppo. Questo non è strano: in effetti la quantità di moto dovrebbe essere una grandezza conservata in assenza di forze esterne, ma se si applicano le trasformazioni di Lorentz potrebbe accadere che la quantità di moto totale di un certo numero di particelle sia conservata in certi sistemi di riferimento e non in altri. La teoria della relatività suggerisce che non sia $m_0 \mathbf{u}$ a restare costante nel tempo in assenza di forze, ma $\underline{p}$. Per velocità basse $\gamma \simeq 1$ e si ottiene per la componente spaziale l'usuale definizione di quantità di moto. Ma quando γ diventa grande la quantità di moto si deve sostituire con $\gamma m_0 \mathbf{u}$ che può anche diventare infinita. È **come se** (e sottolineiamo **come se**) la massa della particella m_0 aumentasse di un fattore γ diventando $m = \gamma m_0$.

Va detto che la massa di una particella è costante anche in relatività. L'affermazione secondo la quale la massa aumenterebbe all'aumentare della velocità è di per sé falsa, anche se in certi casi pensare in questi termini **funziona**! In effetti funziona nei casi in cui ci si ostina a interpretare i fenomeni alla luce dei risultati della fisica classica. Per accelerare una particella e farne variare la quantità di moto, secondo la meccanica classica occorre una forza (per la Legge di Newton $\mathbf{F} = \Delta \mathbf{p} / \Delta t$, quindi $\Delta \mathbf{p} = \mathbf{F} \Delta t$). Man mano che la particella accelera e cambia la sua velocità, se al crescere della velocità la massa crescesse, la forza necessaria per far aumentare anche solo di pochissimo la sua velocità diventerebbe sempre più alta (in questo schema possiamo riscrivere la Legge di Newton come $\mathbf{a} = \mathbf{F} / m$, con m variabile) e così alla fine non si riesce mai a raggiungere la velocità limite della luce. Infatti, secondo questa interpretazione, la grandezza m (che è il rapporto tra la forza e l'accelerazione di un punto materiale, quindi una grandezza fisica **diversa** da quella che si misura con una bilancia) può diventare infinita e, sempre secon-

⁸La quantità di moto di una particella si può chiamare anche *impulso*, anche se questo termine non è del tutto esatto; si riferisce infatti non alla quantità di moto della particella, ma all'integrale nel tempo delle forze agenti sulla particella che numericamente è uguale alla quantità di moto, ma concettualmente è qualcosa di diverso.

do la fisica classica, nessuna forza potrà mai modificarne la velocità. Il fatto è che quando si applica una forza a una particella, accelerandola, non se ne modifica la velocità, come in meccanica classica, ma l'energia! Consigliamo dunque di non assumere mai questa interpretazione e ne parliamo perché si trova su molti libri e su molte pubblicazioni.

Esercizio 11.3 *Elettroni accelerati*

Un fascio di N elettroni è accelerato da una differenza di potenziale di 100 000 V per poi, dopo aver percorso 150 m, andare a collidere con un calorimetro con 2 ℓ d'acqua la cui temperatura s'innalza di 10^{-4} °C. Calcola quanti elettroni sono presenti nel fascio e trova il tempo impiegato a percorrere la distanza che li separa dal calorimetro, una volta usciti dall'acceleratore.

SOLUZIONE →

Uno dei casi in cui l'interpretazione sembra funzionare è quello del calcolo del **raggio di curvatura** di una particella carica in campo magnetico. Il raggio di curvatura aumenta di un fattore γ rispetto a quanto previsto dalla **Forza di Lorentz** e questo potrebbe far pensare al fatto che in effetti la massa aumenta. Va detto che il calcolo della forza di Lorentz in questi casi non è così banale come può sembrare perché la particella carica non si muove affatto di moto rettilineo uniforme rispetto a noi fermi nel laboratorio. Tuttavia se si esegue il calcolo corretto si trova il risultato dell'equazione (16.1). In realtà la questione si può reinterpretare in altri termini: il raggio di curvatura R determina la lunghezza della traiettoria $L = 2\pi R$. Se R è abbastanza grande il moto si può considerare approssimativamente rettilineo e la lunghezza *vista* dalla particella carica è contratta di un fattore γ rispetto a quella misurata da chi si trova fermo nel laboratorio. Poiché la particella carica, nel suo sistema di riferimento, vede una lunghezza pari a quella prevista dalla formula classica di Lorentz, noi, nel laboratorio, la vediamo più lunga di un fattore γ . Di conseguenza il raggio di curvatura aumenta dello stesso fattore.

Come si vede più sotto, l'energia di una particella dipende dalla sua quantità di moto e dalla sua massa. La massa di una particella è un invariante relativistico dato, a meno di una costante c^2 , dalla radice del modulo quadro del quadripulso (eq. (11.50)). La massa effettiva misurata classicamente no: questa cresce con la velocità della particella di un fattore γ per dare una *massa classica* $m = \gamma m_0$ dove m_0 , detta **massa a riposo** della particella, è la massa della particella misurata in un sistema di riferimento in cui è ferma.

Come dobbiamo interpretare, invece, la parte temporale di questo quadripulso, che vale $\gamma m_0 c$? Per capire di che si tratta è utile riscriverlo nell'approssimazione classica, cioè per $v \ll c$. In questo caso

$$\gamma = \frac{1}{\sqrt{1 - \beta^2}} \simeq 1 + \frac{1}{2}\beta^2 \quad (11.45)$$

come si evince dal Paragrafo **approssimazione di funzioni** dell'Appendice, equazione (25.37). La quantità $\gamma m_0 c$ in approssimazione non relativistica diventa quindi

$$\gamma m_0 c \simeq \left(1 + \frac{1}{2}\beta^2\right) m_0 c = m_0 c + m_0 \frac{u^2}{2c} \quad (11.46)$$

che moltiplicata per c dà

$$E = \gamma m_0 c^2 \simeq m_0 c^2 + m_0 \frac{u^2}{2}. \quad (11.47)$$

Il secondo addendo è l'energia cinetica della particella. Evidentemente anche il primo addendo ha le dimensioni di un'energia. Dobbiamo quindi interpretare la componente temporale di questo quadripulso, come l'energia della particella divisa per la costante c . Ma una particella ferma in assenza di forze non possiede energia, mentre in questo caso avremmo che l'energia di una particella in quiete non è nulla, ma vale $m_0 c^2$. Questa grandezza, detta **energia a riposo** della particella, rappresenta dunque l'energia posseduta da una particella ferma per il solo fatto di avere una massa m_0 .

In altre parole potremmo concluderne che la massa di una particella non è altro che un'altra forma di

energia, dal momento che contribuisce a quest'ultima per una quantità proporzionale a essa. Possiamo quindi definire

$$\underline{P} = \left(\frac{E}{c}, \mathbf{p} \right) \quad (11.48)$$

come il quadrivettore **energia-impulso** o **quadrimpulso** che ha la componente temporale pari all'energia della particella misurata in unità di c e quella spaziale pari alla sua quantità di moto (ricordando di usare, per questa, la massa relativistica $m = \gamma m_0$).

Ricaviamo alcune relazioni utili da questo quadrivettore: innanzi tutto conoscendo l'energia totale $E = \gamma m_0 c^2$ di una particella e la sua massa, se ne può ricavare il fattore di Lorentz $\gamma = E/m_0 c^2$ (che in unità naturali diventa $\gamma = E/m_0$). Calcoliamo quindi il modulo quadro del quadrimpulso:

$$P^2 = \frac{E^2}{c^2} - p^2 = \frac{\gamma^2 m_0^2 c^4}{c^2} - \gamma^2 m_0^2 \beta^2 c^2 = \gamma^2 m_0^2 c^2 (1 - \beta^2) \quad (11.49)$$

Osservando che $\gamma^2 = (1 - \beta^2)^{-1}$ si ottiene che

$$P^2 = m_0^2 c^2 \quad (11.50)$$

cioè che il modulo quadro del quadrimpulso è una costante pari alla massa a riposo della particella moltiplicata per c (in unità naturali è proprio uguale alla massa della particella). Si tratta di una relazione molto utile che permette di calcolare lo stato di una particella dopo aver interagito con altre particelle, perché si può imporre che il quadrimpulso totale si conservi e che il suo modulo quadro sia uguale alla massa quadra delle particelle corrispondenti (che è invariante e quindi non dipende dal sistema di riferimento in cui si eseguono le misure).

Esercizio 11.4 Le dimensioni di LHC

I protoni in LHC con un'energia di 7 TeV sono mantenuti nella loro traiettoria da un campo magnetico uniforme di 8.4 T, perpendicolare alla loro velocità. Calcola la lunghezza dell'acceleratore.

SOLUZIONE →

Osserviamo infine che il rapporto p/E tra il modulo della quantità di moto e l'energia totale della particella vale

$$\frac{p}{E} = \frac{\gamma m_0 \beta c}{\gamma m_0 c^2} = \frac{\beta}{c} \quad (11.51)$$

che in unità naturali si riscrive come $p/E = \beta$. È anche utile osservare che, in unità naturali, l'energia ha le stesse dimensioni fisiche della quantità di moto (le componenti di un quadrivettore devono sempre avere le stesse dimensioni fisiche) e che anche la massa ha le stesse dimensioni fisiche dell'energia. Le masse perciò si possono misurare in unità di energia e per trovare il loro valore in unità SI basta dividerne il valore espresso in unità di energia per c^2 . Ad esempio la massa di un protone in unità naturali vale circa 1 GeV. Ricordiamo che $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$, quindi nel SI la massa del protone vale

$$m_p \simeq 1 \times 10^9 \times 1.6 \times 10^{-19} \simeq 1.6 \times 10^{-10} \text{ J}. \quad (11.52)$$

Dividendo per $c^2 = 9 \times 10^{16} \text{ m}^2 \text{s}^{-2}$ si ottiene il valore della massa in kg: $m_p \simeq 1.6 \times 10^{-10} / (9 \times 10^{16}) \simeq 1.8 \times 10^{-27} \text{ kg}$.

La relazione relativistica secondo cui $E^2 - p^2 = m^2$ in unità naturali (in unità SI la stessa relazione diventa $E^2 - p^2 c^2 = m^2 c^4$) ha un'interessante conseguenza: se $m = 0$, $E = p$ e quindi $p/E = \beta = 1$: la particella si muove alla velocità della luce.

Quest'osservazione permette di affermare che, qualora la luce si possa considerare un flusso di particelle, queste devono avere massa nulla.

Energia dalle stelle

Da dove proviene l'energia che emanano le Stelle, che giunge fino a noi sotto forma di luce e calore? In effetti la risposta sta nella fisica relativistica. Nella parte piú interna delle stelle, a causa delle alte pressioni provocate dalla forza gravitazionale, singoli protoni si avvicinano abbastanza da fondere in nuclei di elio combinandosi con neutroni, a loro volta derivanti dalla trasformazione di protoni in neutroni, elettroni e neutrini. La massa di un protone in unità di massa atomica (l'unità corrisponde a $1/12$ della massa del nucleo di ^{12}C) è 1.008, mentre quella di un nucleo di elio è di 4.002602. Partendo da quattro protoni, la cui massa totale è $4 \times 1.008 = 4.032$, si finisce con un nucleo di elio di massa piú bassa con una differenza di poco meno di 0.030. Questa differenza di massa è dissipata sotto forma di energia (cinetica dei prodotti delle reazioni e dei fotoni emessi nel corso del processo). Ogni quattro protoni che fondono dunque il Sole produce circa 30 MeV di energia (un'unità di massa atomica corrisponde circa alla massa di un protone che pesa approssimativamente 1 GeV in unità naturali).

Muoversi tra sistemi

PREREQUISITI: Relatività ristretta.

Secondo quanto abbiamo appreso, per conoscere lo stato di un oggetto non basta conoscerne posizione e velocità: è necessario conoscere anche l'istante di tempo al quale lo stato è riferito e in quale riferimento è espresso. Solo così è possibile esprimere lo stato dello stesso oggetto in qualunque altro sistema di riferimento.

Per capire come comportarsi quando si passa da un sistema di riferimento a un altro possiamo procedere con un esempio, tratto (ma un po' modificato) dalla prima simulazione della prova d'esame di fisica fornita dal Ministero della Pubblica Istruzione nel 2015. Nel problema proposto un'astronave si muove a una velocità pari al 75 % di quella della luce allontanandosi dalla Terra, diretta verso una destinazione molto lontana. Per ragioni tecniche è necessario che dalla Terra sia inviato un segnale radio che interagisca con i sistemi di bordo e questo viene fatto a distanza di un anno dalla partenza dell'astronave. Si chiede quindi di conoscere quanto tempo passa, secondo gli astronauti a bordo, tra il momento della partenza e quello dell'invio del segnale da Terra, nonché il tempo trascorso tra la partenza della navicella e l'arrivo di questo segnale. Nel problema si assume che tutte le accelerazioni siano trascurabili e che il moto sia rettilineo e uniforme.

Come al solito converrà esprimere tutto in unità naturali ($c = 1$). Secondo il controllo a Terra, l'equazione del moto dell'astronave è

$$x = vt \quad (12.1)$$

dove $v = 0.75$ è la velocità dell'astronave misurata a Terra in unità di c , t è il tempo trascorso a Terra

e x la distanza a cui si trova l'astronave al tempo t misurata nel sistema di riferimento terrestre. Dopo un anno, dalla Terra parte un segnale radio, che viaggia a velocità c , la cui equazione del moto si può scrivere come

$$x_S = c(t - t_0), \quad (12.2)$$

dove $t_0 = 1$. Evidentemente il segnale radio raggiunge l'astronave quando $x = x_S$, cioè per

$$vt = c(t - t_0) \quad (12.3)$$

il che avviene per

$$t = \frac{ct_0}{c - v}. \quad (12.4)$$

In quest'istante l'astronave si trova in

$$x = vt = v \frac{ct_0}{c - v}. \quad (12.5)$$

Numericamente abbiamo che

$$\begin{aligned} t &= \frac{1}{1 - 0.75} = 4 \\ x &= 0.75 \frac{1}{1 - 0.75} = 3. \end{aligned} \quad (12.6)$$

Le coordinate quadridimensionali (o, se preferite, il suo quadrivettore spazio-temporale) dell'**evento** che consiste nel raggiungimento dell'astronave da parte del segnale radio sono, nel sistema di riferimento della Terra, $(4, 3)$ dove il primo numero rappresenta il tempo (misurato in anni) e il secondo la distanza (misurata in anni-luce).

Come si ottengono le coordinate dello stesso evento nel sistema di riferimento degli astronauti? È semplice: basta applicare le [trasformazioni di Lorentz](#), secondo le quali

$$\begin{aligned} t' &= \gamma(t - vx) \\ x' &= \gamma(x - vt) \end{aligned} \quad (12.7)$$

dove

$$\gamma = \frac{1}{1 - v^2} \simeq 1.5. \quad (12.8)$$

Il risultato è che l'evento si verifica alle coordinate

$$\begin{aligned} t' &\simeq 1.5(4 - 0.75 \times 3) = 2.63 \\ x' &\simeq 1.5(3 - 0.75 \times 4) = 0. \end{aligned} \quad (12.9)$$

Il fatto che $x' = 0$ è una conferma che il conto fatto è corretto: infatti, nel sistema di riferimento degli astronauti, la coordinata spaziale alla quale il segnale giunge presso la navicella coincide con l'origine del sistema di riferimento, quindi doveva essere per forza $x' = 0$! Le nuove coordinate, nel nuovo sistema di riferimento, sono dunque $(2.63, 0)$.

Per calcolare l'istante di tempo al quale, secondo gli astronauti, il segnale parte dalla Terra è sufficiente applicare la stessa tecnica all'evento di coordinate $(1, 0)$, perché il segnale parte a $t = 1$ dal punto $x = 0$:

$$\begin{aligned} t' &= \gamma t \simeq 1.5 \\ x' &= v\gamma t \simeq -0.75 \times 1.5 \times 1 = -1.125. \end{aligned} \quad (12.10)$$

Questo significa che, secondo gli astronauti, il segnale parte dalla Terra dopo un anno e mezzo dalla loro partenza, quando la Terra si trova a una distanza pari a 1.125 anni-luce da loro (il segno $-$ c'è perché la Terra si sta allontanando dalla parte delle x negative rispetto all'astronave).

Osserviamo che il modulo quadro del quadriettore è invariante per trasformazioni del sistema di riferimento: infatti, nel sistema della Terra

$$s^2 = t^2 - x^2 = 16 - 9 = 7, \quad (12.11)$$

mentre nel sistema di riferimento degli astronauti avremo

$$s'^2 = t'^2 - x'^2 \simeq 6.9 \quad (12.12)$$

che con le approssimazioni fatte coincide con s^2 . In sostanza, ogni volta che si ha la necessità di calcolare il valore di una grandezza fisica misurata in un sistema di riferimento in moto rispetto a un altro, è sufficiente calcolare le coordinate dell'**evento** nel sistema nel quale si hanno i dati, per poi trasformarle con Lorentz.

12.1 Una tecnica alternativa

Un altro modo (più geometrico e grafico) di vedere la stessa cosa è il seguente. Nel piano (t, x) tutti i punti sull'asse delle ascisse sono punti per i quali $x = 0$, mentre quelli sull'asse delle ordinate hanno in comune il valore di $t = 0$. Un **evento** si rappresenta su questo piano come un punto di coordinate (t, v) e il piano in questione prende il nome di **piano (o spazio) di Minkoski**¹.

Tutti gli eventi **simultanei**, che cioè avvengono nello stesso istante di tempo, sono quelli che si trovano lungo rette verticali, parallele all'asse spaziale x . Al contrario, tutti gli eventi che avvengono nel stesso punto (ma a istanti diversi) sono disposti su rette parallele all'asse delle ascisse.

Trasformando con Lorentz un punto sull'asse dei tempi di un sistema in moto con velocità v rispetto a quello considerato ($x' = 0$) si trova che

$$\begin{aligned} x &= v\gamma t' \\ t &= \gamma t'. \end{aligned} \quad (12.13)$$

Dividendo membro a membro si trova che $\frac{x}{t} = v$ o $x = vt$ che, nel piano di Minkoski, si rappresenta come una retta passante per l'origine di coefficiente angolare v . Tutti i punti di questa retta hanno in comune il fatto che in ciascuno di essi $x' = 0$: sono quindi punti dell'asse temporale nel sistema in moto: rappresentano quindi l'asse delle t' nel sistema in moto. I punti che invece sono simultanei nel sistema

¹In molti libri il piano di Minkoski ha la coordinata temporale lungo l'ordinata e quella spaziale lungo l'ascissa: evidentemente è la stessa cosa, ma usando questa convenzione si rischia di commettere errori grossolani dovuti al fatto che siamo abituati a vedere le equazioni del moto rappresentate su un piano (t, x) .

in moto e per i quali $t' = 0$, trasformati con Lorentz nel sistema fermo, hanno coordinate

$$\begin{aligned} x &= \gamma x' \\ t &= v\gamma x'. \end{aligned} \quad (12.14)$$

Dividendo ancora membro a membro si trova $x = \frac{t}{v}$, che nel piano di Minkowski è una retta passante per l'origine, di pendenza $1/v$ e simmetrica, rispetto alla bisettrice del primo quadrante, a quella con pendenza v . Su questa retta giacciono tutti i punti che nel sistema in moto hanno coordinata $t' = 0$ e quindi questa retta rappresenta l'asse delle x' in questo sistema.

In definitiva, la trasformazione di Lorentz ha l'effetto di *schiacciare* gli assi verso la bisettrice del primo quadrante (vedi Figura 12.1). In effetti gli assi t' e x' coincidono entrambi con la bisettrice quando $v = c$. È come se gli assi dello spazio di Minkowski si potessero rappresentare usando righelli rigidi *incernierati* agli estremi: schiacciando con le mani due vertici opposti del quadrato che rappresenta il sistema fermo, la forma assunta dai righelli cambia, assumendo quella di un rombo. Maggiore è la velocità relativa tra i sistemi e più è *schiacciato* il sistema in moto. È da notare che, in ciascuna delle due situazioni, la lunghezza dell'unità di misura non cambia, perciò chiunque esegua una misura di distanza con questi righelli ha la sensazione che nulla cambi con il variare dello stato di moto. Ma se si confrontano le misure fatte dall'uno e dall'altro, queste non coincidono (basta osservare che i punti in cui s'incrociano le rette della griglia nera non coincidono più con quelli in cui s'incrociano le rette della griglia rossa).

Gli eventi che nel sistema in moto avvengono simultaneamente si trovano tutti su rette parallele all'asse x' , mentre gli eventi che avvengono a tempi diversi, ma nello stesso punto, sono allineati su rette parallele all'asse t' .

Per disegnare il grafico in questione puoi procedere in questo modo: il punto che ha coordinate $(1, 0)$ nel sistema in moto (in rosso nella figura), in quello fermo ha coordinate $(\cos \theta, \sin \theta)$, dove θ è l'angolo formato tra l'asse t e quello t' , che si ricava come

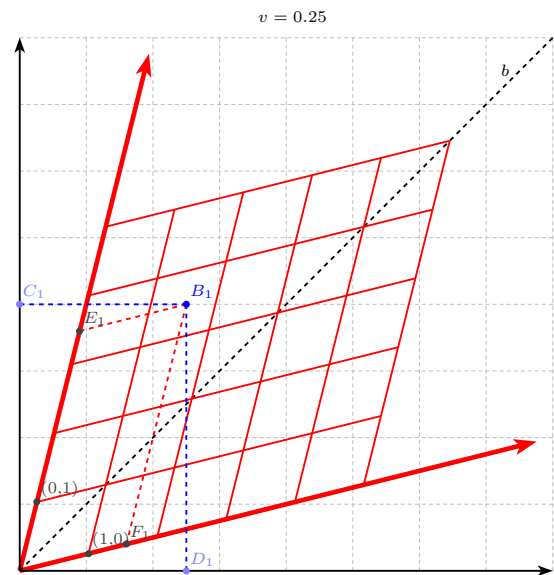


Figura 12.1 Nel sistema di riferimento rappresentato in nero il punto B_1 ha coordinate (D_1, C_1) . Una trasformazione di Lorentz in cui $v = 0.25$ cambia le coordinate di B_1 in (F_1, E_1) . Puoi sperimentare l'effetto che fa la trasformazione di Lorentz al variare di v all'indirizzo <http://tube.geogebra.org/student/m836297>.

$\theta = \arctan v$, essendo v la pendenza della retta che giace sull'asse x' . Analogamente le coordinate del punto $(0, 1)$ del sistema in moto sono $(\sin \theta, \cos \theta)$. Individuati questi due punti sul piano basta tracciare le rette passanti per l'origine e per questi due punti per avere la direzione degli assi coordinati nel sistema in moto. Le coordinate spazio-temporali di un qualunque evento B_1 nel piano si ricavano nel sistema fermo tracciando le rette parallele agli assi x e t e trovando i punti d'intersezione con gli assi di questo sistema (in nero). Quelle nel sistema in moto si ricavano conducendo da B_1 le rette parallele agli assi x' e t' . I punti in cui queste rette intersecano gli assi x' e t' , misurati in unità di distanza tra l'origine e i punti di coordinate $(0, 1)$ e $(1, 0)$ nel sistema in moto, rappresentano le coordinate dell'evento nel

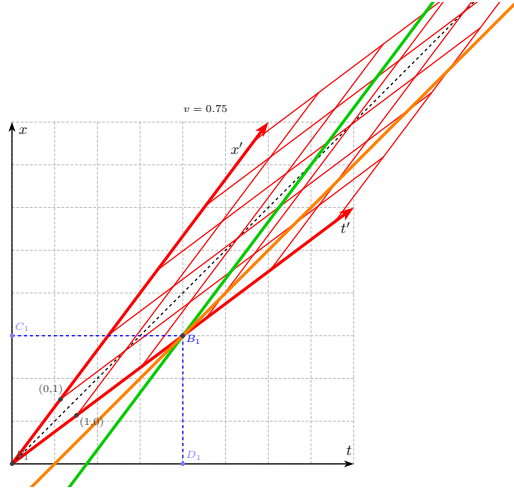


Figura 12.2 L'equazione del moto dell'astronave, $x = vt$, si rappresenta come una retta coincidente con l'asse t' . Il segnale radio partito dalla Terra dopo un anno è rappresentato dalla linea arancione. Il punto d'intersezione è B_1 che rappresenta l'evento consistente nel raggiungimento dell'astronave da parte del segnale. Osservate che la coordinata spaziale dell'evento coincide con la partenza del segnale radio per gli astronauti è negativa.

nuovo sistema.

Un esempio chiarirà il meccanismo: secondo i tecnici rimasti a Terra dell'esempio riportato al paragrafo precedente, le coordinate spazio-temporali del punto in cui il segnale radio inviato da Terra dopo un anno dalla partenza raggiunge l'astronave, sono $B_1 = (4, 3)$ (Figura 12.2), dove il primo numero rappresenta il tempo t misurato in anni dai tecnici a Terra e il secondo la distanza x alla quale si trova l'astronave secondo costoro.

Le rette che rappresentano gli assi coordinati nel sistema in moto hanno equazione $x = vt$ e $x = \frac{t}{v}$, rispettivamente. La retta parallela all'asse t' passante per B_1 ha equazione

$$x = vt + q \quad (12.15)$$

e il valore di q si trova imponendo il passaggio per B_1 :

$$3 = 0.75 \times 4 + q \quad (12.16)$$

da cui si ricava che $q = 0$. Questo risultato è in accordo con quello che abbiamo trovato prima, perché implica che l'evento si trovi sull'asse delle t' , dovendo avere $x' = 0$. Anche graficamente si vede che il punto B_1 sta sull'asse t' . La retta parallela all'altro asse invece ha equazione

$$x = \frac{t}{v} + q \quad (12.17)$$

e di nuovo q si determina imponendo il passaggio per B_1 :

$$3 = \frac{4}{0.75} + q \quad (12.18)$$

per cui $q \simeq -2.33$. Cerchiamo l'intersezione tra questa retta (rappresentata in verde nella figura) e l'asse delle t' : dobbiamo mettere a sistema l'equazione di questa retta con quella dell'asse per cui avremo che

$$\frac{t}{v} + q = vt \quad (12.19)$$

il che si verifica quando

$$t \left(v - \frac{1}{v} \right) = q \quad (12.20)$$

cioè per

$$t = \frac{qv}{v^2 - 1} = -qv\gamma \simeq 2.62. \quad (12.21)$$

Potete verificare che il risultato è corretto osservando che il punto B_1 nella Figura 12.2 si trova circa a metà tra il punto di coordinate $(2, 0)$ e quello di coordinate $(3, 0)$ nel sistema (t', x') . Algebricamente avevamo trovato lo stesso risultato (a meno di un decimale, a causa degli arrotondamenti fatti).

Le stesse coordinate di B_1 si possono facilmente trovare con questo metodo grafico: l'equazione del moto dell'astronave, per chi sta a Terra, è infatti $x = vt$ e quindi la sua rappresentazione grafica coincide con l'asse delle t' . Quella del segnale radio è

$x = c(t - t_0)$ con $t_0 = 1$, che nella rappresentazione grafica è una retta passante di pendenza $c = 1$ che passa per il punto $(t_0, 0)$, che è poi quella riprodotta in arancio nella figura. Il punto in cui il segnale raggiunge l'astronave è quello in cui la retta arancione e quella di equazione $x = vt$ (in rosso, coincidente con l'asse t') s'incrociano: si vede subito che questo punto è B_1 .

Il metodo grafico può essere un po' più laborioso di quello algebrico, ma è un sistema rapido per verificare, almeno a occhio, di non aver commesso errori con l'algebra. Basta un righello e un foglio di carta: si tracciano gli assi coordinati del sistema in moto e si rappresentano le equazioni del moto degli oggetti nel piano (t, x) . I segnali luminosi (e tutti quelli che in generale si muovono a velocità c) si rappresentano come rette parallele alla bisettrice, mentre per trovare le coordinate dei punti nel sistema in moto basta condurre le rette parallele ai nuovi assi passanti per gli eventi d'interesse.

12.2 Acceleratori e collider

Un altro effetto della trasformazione di una grandezza fisica da un sistema di riferimento a un altro in moto relativo si ha negli [esperimenti di fisica delle particelle](#), nei quali si produce lo scontro tra una particella proiettile e una bersaglio.

Dal momento che l'energia e la massa sono di fatto la stessa cosa possiamo pensare di produrre energia *consumando* massa o viceversa. La produzione di energia a spese della massa è il meccanismo che tiene accese le stelle (vedi box) e le centrali nucleari nelle quali i nuclei di uranio sono spezzati in parti più piccole la cui somma delle masse è inferiore a quella del nucleo originario.

Per creare massa dall'energia si può far urtare una particella di energia abbastanza elevata su un bersaglio: l'energia della collisione è disponibile per la produzione di particelle. Per capire quanta energia serve usiamo le proprietà dei quadrivettori, considerando dapprima il caso di una particella con massa m , di energia E e quantità di moto $\mathbf{p}$ che urta contro un bersaglio costituito di particelle dello stesso tipo

(ad esempio, un fascio di protoni contro un bersaglio di carbonio, a sua volta costituito di protoni). Il quadrimpulso totale di questo sistema è

$$(E, \mathbf{p}) + (m, \mathbf{0}) = (E + m, \mathbf{p}) . \quad (12.22)$$

Il modulo quadro di questo quadrivettore è

$$s^2 = E^2 + m^2 + 2Em - p^2 \quad (12.23)$$

ma $E^2 - p^2 = m^2$ e quindi

$$s^2 = 2m(E + m) . \quad (12.24)$$

Se nell'urto si potesse produrre una particella di massa M questa avrebbe una massa pari a $\sqrt{s^2}$, perché il suo quadrimpulso sarebbe, per la conservazione di quest'ultimo, pari alla somma dei quadrimpulsi delle particelle iniziali. Dunque per produrre particelle di massa M occorrono fasci di protoni di energia pari a

$$E = \frac{M^2}{2m} - m . \quad (12.25)$$

Possiamo anche dire che l'energia disponibile per la produzione di massa cresce come la radice dell'energia del fascio. Se volessimo produrre un [bosone di Higgs](#) usando fasci di protoni, sapendo che $m \simeq 1$ GeV e $M \simeq 125$ GeV abbiamo

$$E \simeq \frac{125^2}{2} - 1 \simeq 7800 \text{ GeV} . \quad (12.26)$$

È una quantità d'energia enorme (trasformatela in Joule per rendervene conto). Se invece volessimo produrre questa stessa particella usando fasci di uguale energia che collidono l'uno contro l'altro il quadrimpulso totale sarebbe

$$(E, \mathbf{p}) + (E, -\mathbf{p}) = (2E, \mathbf{0}) . \quad (12.27)$$

In questo caso l'energia disponibile per produrre una particella è

$$s = \sqrt{s^2} = \sqrt{4E^2} = 2E , \quad (12.28)$$

cioè la semplice somma delle energie delle particelle collidenti. Per produrre un bosone di Higgs dunque

basterebbero due protoni di energia pari a $125/2 = 62.5$ GeV. Molto meno di prima!²

È per questo che si costruiscono gli acceleratori circolari **collisori** o **collider**: l'energia disponibile per produrre nuove particelle infatti cresce linearmente con la massa delle particelle da produrre, mentre nel caso di macchine a bersaglio fisso cresce come il quadrato della massa.

²In realtà occorre molto di più perché i protoni non sono particelle puntiformi, ma composte di *quark*. A scontrarsi dunque sono questi ultimi, più leggeri e con energia minore.

La Relatività Generale

PREREQUISITI: trasformazioni di coordinate, sistemi inerziali e non inerziali, dinamica elementare, relatività ristretta

La teoria della relatività ristretta permette di predire il comportamento osservato di oggetti che si muovono di moto rettilineo uniforme rispetto a un osservatore. Appare subito chiaro che limitare le predizioni a questo tipo di moto non è per nulla soddisfacente: nulla si muove realmente di moto rettilineo uniforme. È dunque opportuno studiare i moti relativi di sistemi accelerati (non inerziali), alla luce dei risultati ottenuti avendo discusso la relatività ristretta. È questo l'oggetto della teoria della **relatività generale**.

Sfortunatamente l'apparato matematico necessario a trattare questo tipo di moti è particolarmente complesso e dunque risulta sempre molto complicato introdurre questa teoria ai non esperti. In questo capitolo proviamo quanto meno a introdurre le linee generali e i concetti alla base di questa teoria.

13.1 La misura nei vari sistemi di riferimento

La [relatività ristretta](#) insegna che per definire lo stato di un evento in un sistema di riferimento occorre fornire oltre alla posizione degli oggetti che si stanno misurando, anche il tempo della misura misurato da un orologio solidale con il sistema di riferimento usato per misurare le posizioni. Secondo questa teoria, passando da un sistema di riferimento a un altro in moto rettilineo uniforme rispetto al primo, le coordinate dei quadrivettori si trasformano con le

trasformazioni di Lorentz, che possiamo scrivere in forma compatta come

$$x'^{\mu} = L^{\mu}_{\nu} x^{\nu} \quad (13.1)$$

dove nell'espressione sopra riportata s'intende la somma su tutti gli indici ripetuti e $\mu = 0, \dots, 3$ rappresentano gli indici relativi alle componenti del quadrivettore. In sostanza $\underline{x} = (t, x, y, z) = (x^0, x^1, x^2, x^3)$. L^{μ}_{ν} è qualcosa che ha due indici, le cui componenti rappresentano i coefficienti della trasformazione. Applicando l'equazione (13.1) si ha quindi che

$$x'^0 = L^0_0 x^0 + L^0_1 x^1 + L^0_2 x^2 + L^0_3 x^3 \quad (13.2)$$

e nello specifico, sapendo che $t' = \gamma(t - \beta x)$, abbiamo che $L^0_0 = \gamma$, $L^0_1 = -\gamma\beta$, $L^0_2 = 0$ e $L^0_3 = 0$. Allo stesso modo, dalla trasformazione $x' = \gamma(x - \beta t)$ e dall'equazione

$$x'^1 = L^1_0 x^0 + L^1_1 x^1 + L^1_2 x^2 + L^1_3 x^3 \quad (13.3)$$

si ricava che $L^1_0 = -\gamma\beta$, $L^1_1 = \gamma$, $L^1_2 = 0$ e $L^1_3 = 0$. Non è difficile trovare tutte le 16 componenti di L , che si può esprimere come una matrice. È abbastanza chiaro che fare una teoria che permetta di scrivere le equazioni del moto in sistemi di riferimento qualunque significa fare una teoria che consenta di individuare le componenti della matrice L (che più propriamente si chiama **tensore**, che in matematica definisce le relazioni geometriche esistenti tra le coordinate dei vettori in uno spazio qualunque).

Possiamo usare un tensore per esprimere la lunghezza del quadrivettore, definendo il modulo quadro di un quadrivettore $\underline{x}$, che sappiamo essere un'invariante, come

$$s^2 = g_{\mu\nu} x^\mu x^\nu. \quad (13.4)$$

Ricordando che questa scrittura significa che occorre sommare sugli indici ripetuti abbiamo che

$$s^2 = g_{\mu 0} x^\mu x^0 + g_{\mu 1} x^\mu x^1 + g_{\mu 2} x^\mu x^2 + g_{\mu 3} x^\mu x^3 \quad (13.5)$$

in cui ancora l'indice μ è ripetuto e perciò

$$\begin{aligned} s^2 = & g_{00} x^0 x^0 + g_{10} x^1 x^0 + g_{20} x^2 x^0 + g_{30} x^3 x^0 + \\ & g_{01} x^0 x^1 + g_{11} x^1 x^1 + g_{21} x^2 x^1 + g_{31} x^3 x^1 + \\ & g_{02} x^0 x^2 + g_{12} x^1 x^2 + g_{22} x^2 x^2 + g_{32} x^3 x^2 + \\ & g_{03} x^0 x^3 + g_{13} x^1 x^3 + g_{23} x^2 x^3 + g_{33} x^3 x^3 \end{aligned} \quad (13.6)$$

Confrontando con l'espressione della lunghezza dell'invariante di Lorentz $s^2 = t^2 - (x^2 + y^2 + z^2) = (x^0)^2 - ((x^1)^2 + (x^2)^2 + (x^3)^2)$ si vede subito che $g_{\mu\nu} = 0$ ogni volta che $\mu \neq \nu$ e in caso contrario vale -1 se $\mu = \nu \neq 0$ e $+1$ se $\mu = \nu = 0$.

Il tensore $g_{\mu\nu}$ si chiama **tensore metrico** perché serve per calcolare la lunghezza dei segmenti che hanno un estremo nell'origine del sistema di riferimento e l'altro nel punto scelto. I valori di $g_{\mu\nu}$ dipendono, naturalmente, dalle caratteristiche del sistema di riferimento scelto per esprimere le coordinate e dalle proprietà dello spazio. Ce ne possiamo rendere conto facilmente considerando un sistema di riferimento cartesiano bidimensionale, in cui il tensore metrico $g_{\mu\nu}$ è una matrice 2×2 . In questo sistema il punto $(x, y) = (x^0, x^1)$ individua un punto che dista dall'origine una quantità r pari a $r^2 = (x^0)^2 + (x^1)^2$ cosicché in questo sistema di riferimento $g_{00} = g_{11} = +1$ e $g_{01} = g_{10} = 0$. Se invece scegliamo di rappresentare il punto in coordinate cartesiane abbiamo che le coordinate sono $(r, \theta) = (x^0, x^1)$ dove

$$\begin{cases} x = x^0 = r \cos \theta = x'^0 \cos x'^1 \\ y = x^1 = r \sin \theta = x'^0 \sin x'^1. \end{cases} \quad (13.7)$$

In questo caso la lunghezza quadra del vettore è semplicemente $r^2 = (x^0)^2$, quindi $g_{\mu\nu} = 0$ sempre tranne che per $\mu = \nu = 0$ per il quale vale $g_{00} = 1$.

In entrambi i casi il valore numerico della distanza sarebbe lo stesso, ma in uno spazio **curvo** i valori di $g_{\mu\nu}$ sono tali da cambiare la distanza tra due punti rispetto a quella che si avrebbe in uno spazio piatto. Se calcoliamo la distanza tra due punti su un piano otteniamo un valore che è diverso da quello che si otterrebbe se questo piano fosse *adagiato* sulla superficie curva della Terra. La distanza euclidea tra i punti è la corda che unisce i due punti adagiati su una sfera, mentre la distanza effettiva sarebbe la lunghezza dell'arco di circonferenza che li unisce. In generale per **spazio curvo** s'intende uno spazio in cui le relazioni geometriche sono tali da non riprodurre i risultati della geometria euclidea.

Quello che ci possiamo aspettare è che i valori assunti dal tensore metrico dipendano dalla trasformazione. Nel caso delle trasformazioni di Lorentz il tensore metrico è simmetrico e sulla diagonale ci sono i valori $(+1, -1, -1, -1)$. Ma se si passa a descrivere la fisica in un sistema di riferimento non inerziale le trasformazioni non saranno più quelle di Lorentz e i valori del tensore metrico cambiano.

Naturalmente nel caso inerziale la trasformazione assume sempre la stessa forma, perché il moto relativo tra i sistemi di riferimento è unico: cambiano solo i valori, non la forma della trasformazione. Nel caso invece dei sistemi di riferimento non inerziali il moto relativo potrebbe essere qualunque e quindi non è possibile scrivere una trasformazione generica come quella di Lorentz che valga per tutti i sistemi. Ogni moto avrà la sua propria trasformazione. Se ne potrebbe trarre la conseguenza che quindi non si potrebbe imparare niente di realmente nuovo da questo studio, ma in realtà non è così.

13.2 Il principio di equivalenza

Già sappiamo che quando si descrive il moto di qualcosa osservandolo da un sistema di riferimento non inerziale, nelle equazioni del moto compaiono dei termini che hanno la stessa forma matematica di una forza e per questo diciamo che nei sistemi di riferimento non inerziali si producono le **forze apparenti**. Queste forze sono dette apparenti proprio

perché non esistono in quanto tali: non sono forze nel senso che non sono il prodotto di un'interazione. Sono solo matematicamente equivalenti a forze.

Un esempio ben noto è la **forza centrifuga**. Se mettiamo in rotazione con le mani un sasso legato a uno spago, sulle dita percepiamo una forza che tira verso il sasso. Molti ritengono erroneamente che questa forza sia la forza centrifuga. In realtà non è altro che la tensione dello spago che si manifesta per effetto della terza Legge di Newton dal momento che all'altro capo è presente una forza **centripeta**, diretta cioè verso la mano. In assenza dello spago il sasso si muoverebbe di moto rettilineo allontanandosi dalla mano. Trascinando con sé lo spago, anche i punti di questo si dovrebbero muovere di moto rettilineo. Ma quando lo spago si tende, le dita della mano lo trattengono applicando una forza diretta verso la mano stessa. Il punto di contatto tra la mano e lo spago non si muove perché si destina una forza di reazione diretta verso l'esterno che annulla l'effetto della forza applicata dalla mano. La forza di reazione è il risultato della somma delle forze (di tipo elettromagnetico) che si esercitano tra le particelle di cui è composto lo spago e che impediscono a questo di dissolversi. Di conseguenza ogni punto dello spago è soggetto a questa forza e a una forza uguale e contraria per effetto della terza Legge di Newton. All'altro capo dello spago (quello cui è legato il sasso) dunque c'è una forza diretta verso la mano prodotta dalle forze interne dello spago, non annullata da altre forze. È questa forza, centripeta appunto, che produce l'accelerazione che fa cambiare la direzione della velocità e fa muovere il sasso di moto circolare.

Se però osserviamo lo stesso fenomeno stando seduti sul sasso, vedremmo il sasso fermo accanto a noi! Se il sasso è fermo e resta tale vuol dire che non ci sono forze in questo sistema e per far valere la seconda Legge di Newton siamo costretti a scrivere che, in questo particolare sistema di riferimento,

$$m\mathbf{a} = \mathbf{F} + \mathbf{F}_{app} \quad (13.8)$$

dove $\mathbf{F}$ è la risultante delle *vere* forze agenti sul sasso (la tensione dello spago), e $\mathbf{F}_{app}$ è qualcosa che

ha le dimensioni di una forza e che deve essere tale da annullare l'accelerazione. Se $\mathbf{F}_{app} = m\mathbf{a}$ vediamo subito che il risultato è che la risultante delle forze applicate è nulla. Nel caso del moto circolare uniforme l'accelerazione è diretta verso il centro della traiettoria e vale v^2/r dove v è il modulo della velocità del sasso e r la sua distanza dall'asse di rotazione. La forza centrifuga quindi vale

$$\mathbf{F}_{app} = m \frac{v^2}{r} \hat{\mathbf{r}}. \quad (13.9)$$

La forza centripeta (che invece è una forza *vera*, effettivamente dovuta a una qualche interazione), vale

$$\mathbf{F} = -m \frac{v^2}{r} \hat{\mathbf{r}}, \quad (13.10)$$

e l'accelerazione del sasso (misurata nel sistema di riferimento in cui è fermo) è nulla. Da quanto sopra se ne deduce che scrivere le trasformazioni che permettono di passare da un sistema di riferimento a un altro in moto accelerato rispetto al primo equivale a trattare un problema in cui sia localmente presente una forza. E non una forza qualsiasi, ma una forza di tipo gravitazionale. Perché sappiamo che la seconda Legge di Newton dice che l'accelerazione $\mathbf{a}$ subita da un corpo di massa m_i è proporzionale alla forza applicata

$$\mathbf{a} = \frac{\mathbf{F}}{m_i}, \quad (13.11)$$

ma al contempo sappiamo che la forza di gravità si scrive $\mathbf{F} = m_G \mathbf{g}$ dove $m_G = m_i$. Quest'ultimo fatto è un fatto sperimentale. Non c'è alcuna ragione per la quale il coefficiente che sta a denominatore nella Legge di Newton debba essere uguale alla costante di accoppiamento delle interazioni gravitazionali. La fisica relativistica dunque dovrà essere equivalente (almeno localmente) alla fisica della gravitazione.

Se non possiamo eseguire misure di oggetti molto distanti da noi (se ad esempio siamo chiusi in una stanza senza finestre) e vediamo degli oggetti cadere non possiamo sapere se questi cadono perché c'è un campo di forze gravitazionali oppure perché l'intera stanza è accelerata verso l'alto. Potremmo saperlo

studiando il moto di oggetti lontani come i pianeti, per questo diciamo che il principio di equivalenza è valido **localmente**.

Da quanto abbiamo detto sopra si capisce che la trasformazione da applicare passando da un sistema a un altro determina le proprietà geometriche dello spazio-tempo e dal momento che trovarsi in un sistema di riferimento accelerato equivale a trovarsi in un campo gravitazionale, ne concludiamo che la presenza di un campo gravitazionale è equivalente alla presenza di una geometria **distorta** dello spazio-tempo nei pressi delle sorgenti del campo (il campo è più intenso vicino ai corpi massicci e lì la geometria somiglierà meno a quella euclidea rispetto a quella che si ha più lontano).

La relatività generale, in definitiva, permette di scrivere in generale la trasformazione di coordinate da un sistema di riferimento a un altro sistema accelerato rispetto al primo, ma oltre a ciò rappresenta una teoria del campo gravitazionale inteso come una distorsione delle proprietà geometriche dello spazio-tempo dovuta alla presenza delle masse, che sono le sorgenti del campo gravitazionale.

13.3 la geometria dell'Universo

Di fatto la relatività generale è una teoria sulla geometria dello spazio (o, per essere più precisi, dello **spazio-tempo**, cioè dello spazio quadridimensionale introdotto con la teoria della relatività ristretta). È interessante osservare che l'idea che lo spazio e il tempo non siano qualcosa di assoluto, ma che la loro *percezione* dipenda dall'osservatore, non è stata introdotta per la prima volta da Einstein con le sue teorie: già il filosofo e matematico **Leibniz** nel 1715 aveva intuito che lo spazio, così come il tempo, andavano considerati come "qualcosa di puramente relativo". Lo spazio, per Leibniz, "è un ordine delle coesistenze", mentre il tempo "è un ordine delle successioni" [11].

La disputa sulla natura dello spazio e del tempo, quindi, è molto antica. Anche **Newton**, nei *Princi-*

La lettera di Leibniz a Clarke

In una delle lettere che Leibniz invia a **Samuel Clarke**, pubblicate dopo la sua morte, si legge: "As for my own opinion, I have said more than once that I hold space to be something purely relative, as time is, that I hold it to be an order of coexistences, as time is an order of successions.". E ancora "I have many demonstrations to confute the fancy of those who take space to be a substance or at least an absolute being." e ne illustra una, di tali confutazioni (che per la scienza moderna suona poco appropriata, ma l'idea di base è in fondo corretta): "without the things placed in it [the space, ndr], one point of space absolutely does not differ in any respect whatsoever from another point of space. Now from this it follows [...] that it is impossible there should be a reason why God, preserving the same situations of bodies among themselves, should have placed them in space after one certain particular manner and not otherwise". In sostanza Leibniz afferma che lo spazio non può essere assoluto perché non ci sarebbe nessuna ragione per la quale gli oggetti che noi osserviamo sono qui e non altrove. Se fossero disposti nello stesso modo in un altro punto di uno spazio assoluto non potremmo distinguere questa situazione da quella attuale, dunque quello che conta non è lo spazio in sé, ma le relazioni tra gli oggetti.

pia [7], nello *scolio*¹ alla fine del capitolo sulle definizioni, cita un esperimento ideale per dimostrare l'assolutezza dello spazio: si prenda un secchio pieno d'acqua sospeso a una corda attorcigliata. Lasciando andare la corda il secchio comincia a ruotare rapidamente, mentre l'acqua, all'inizio, è ferma. Man mano che la corda si svolge, l'acqua si dispone in modo tale da formare una superficie concava. Il moto dell'acqua relativo al secchio non ha alcuna importanza, per Newton, dal momento che rispetto al secchio l'acqua sarebbe ferma. Il fatto che l'acqua

¹Con questo termine s'intendeva una specie di commento conclusivo al termine di una serie di dimostrazioni formali.

risalga i bordi del secchio implica l'esistenza di uno spazio assoluto rispetto al quale l'acqua *sa* di ruotare. **Ernst Mach** confutò questo esperimento (molti anni dopo, alla fine del XIX secolo) osservando che il moto dell'acqua nel secchio deve essere in qualche modo il risultato delle interazioni con il resto dell'Universo: se infatti togliessimo tutto dall'Universo tranne il secchio con l'acqua non potremmo neanche dire che il secchio sta ruotando, perché non si potrebbe stabilire rispetto a cosa stia ruotando.

L'idea alla base della relatività generale di Einstein è che l'acqua nel secchio s'incurva, quando ruota, perché interagisce (gravitazionalmente) con il resto dell'Universo e questa interazione non è altro che il risultato della geometria assunta dall'Universo per il fatto che in esso i corpi assumono certe relazioni spazio-temporali. Un Universo in cui sia presente una sola particella puntiforme è un Universo nel quale lo spazio è semplicemente un concetto privo di senso: lo spazio **non esiste** in questo Universo. Con due particelle si può solo definire la distanza tra le due particelle, quindi al più si può definire uno spazio unidimensionale, ma anche in questo caso avremmo qualche problema, perché non avremmo la possibilità di definire un'unità di misura. Lo spazio più semplice, che abbia un minimo interesse, che possiamo immaginare è formato da almeno tre particelle. Se una delle tre si sposta rispetto alle altre possiamo misurare l'esistenza di un fenomeno, che consiste nella variazione della posizione della particella rispetto alle altre. Lo spazio, dunque, come ente assoluto, non esiste. Lo spazio è semplicemente determinato dalle cose che vi sono immerse: è una relazione tra oggetti.

Se la materia nell'Universo fosse infinita, lo sarebbe anche lo spazio, ma se così non è lo spazio stesso deve essere finito. Quest'affermazione può generare un po' di confusione perché uno s'immagina che recandosi in prossimità della galassia più lontana dal centro dell'Universo dovrebbe vedere una specie di confine oltre il quale non c'è nulla pur tuttavia si potrebbe immaginare l'esistenza di uno spazio.

Un modo per superare l'*impasse* consiste nell'assumere che lo spazio non sia **piatto** come quello euclideo, nel quale due rette parallele non s'incontrano

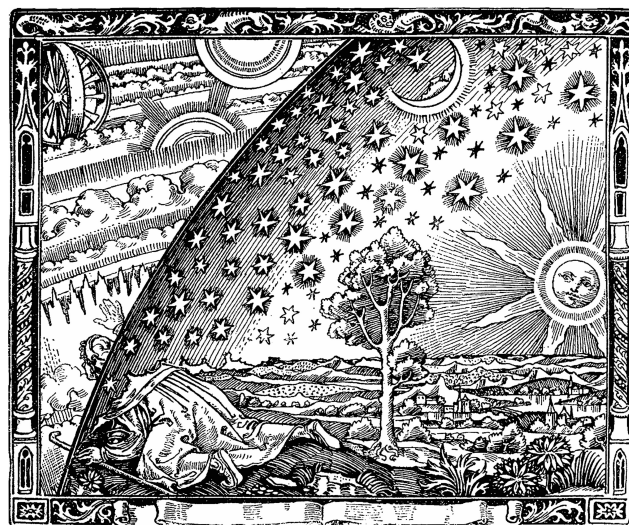


Figura 13.1 Lo spazio è una relazione tra oggetti e non è qualcosa di assoluto. Per questo non esiste un **confine** dello spazio, come quello immaginato da **Camille Flammarion** in quest'incisione.

mai, ma **curvo**. Anche se può sembrare strano, uno spazio curvo non è una cosa così astrusa: supponiamo di essere fermi al Polo Nord in compagnia di un amico e di decidere di separarci procedendo in due direzioni opposte sempre dritti davanti al nostro naso. In uno spazio piatto non c'incontreremmo più con il nostro amico e sarebbe meglio salutarsi per sempre. Ma sulla Terra, sulla quale lo spazio non è piatto, dopo un po' di tempo (e qualche difficoltà di ordine pratico che assumiamo di poter superare) ci vedremmo venire incontro l'un l'altro in prossimità del Polo Sud. Se poi continuassimo a camminare dritti (non prima di esserci salutati di nuovo) prima o poi c'incontreremmo nuovamente al Polo Nord e così via all'infinito. Pur camminando sempre in avanti non raggiungeremmo mai un confine, un limite. Nonostante questo sappiamo bene che la superficie della Terra non è infinita. Ci troviamo, in questo caso, in uno spazio finito, ma illimitato.

La superficie della Terra è uno spazio curvo bidimensionale (ci si può muovere solo su due dimensioni, almeno in prima approssimazione). Se descri-

viamo il nostro moto in uno spazio tri-dimensionale la geometria è euclidea, naturalmente. Non lo è solo quando scriviamo le relazioni tra gli oggetti usando solo due numeri invece che tre.

Se lo spazio tri-dimensionale cui siamo abituati fosse curvo potremmo fare il seguente esperimento: con il nostro amico ci piazziamo in un punto dell'Universo e guardiamo entrambi nella stessa direzione. Uno dei due resta fermo e l'altro comincia a camminare sempre dritto davanti a sé. Cammina, cammina, prima o poi succederà che il nostro amico ce lo vedremo arrivare alle nostre spalle!

In uno spazio come questo le relazioni geometriche tra gli oggetti cambiano, anche se in maniera per noi impercettibile. Resta però valido il principio generale secondo il quale le distanze quadre tra i punti dello spazio e l'origine delle coordinate si possono scrivere come

$$s^2 = g_{\mu\nu} x^\mu x^\nu. \quad (13.12)$$

L'unica differenza consiste nel fatto che ora le componenti di $g_{\mu\nu}$ non sono più quelle di prima, e in generale non sono costanti. Capire come si scrivono le $g_{\mu\nu}$ in funzione della distribuzione delle masse è un compito troppo arduo per un non esperto, che richiede una matematica troppo avanzata. Ma alcune considerazioni di carattere generale possiamo farle.

Dal momento che per la relatività ristretta massa ed energia sono la stessa cosa (nel senso che l'una si può trasformare nell'altra), la geometria dell'Universo non dipenderà solo dalla distribuzione di materia, ma anche dalla distribuzione di energia.

Inoltre, per ogni componente di $g_{\mu\nu}$ si potrebbe scrivere, in linea di principio, un'equazione che dice come varia la componente in esame in funzione della distribuzione di materia ed energia. Si possono cioè scrivere 16 equazioni che ci dicono come variano le diverse componenti di $g_{\mu\nu}$ nei diversi punti dello spazio-tempo in relazione alla presenza, in certi punti, di materia o energia. A partire da queste 16 equazioni si può costruire un oggetto che è anch'esso un tensore, detto **Tensore di Riemann** $R_{\mu\nu}$. In effetti $R_{\mu\nu}$ contiene le **derivate** di $g_{\mu\nu}$ (le sue variazioni) rispetto alle coordinate dello spazio-

tempo. Con questo tensore e con il tensore metrico si può costruire un oggetto che è anch'esso un tensore, perché è la somma di due tensori:

$$G_{\mu\nu} = R_{\mu\nu} + \frac{1}{2}g_{\mu\nu}R \quad (13.13)$$

dove R è una combinazione delle componenti di $R_{\mu\nu}$. Questo tensore non può che essere uguale a un altro tensore, le cui componenti devono dipendere dalla distribuzione di materia ed energia: il **tensore energia-impulso**² $T_{\mu\nu}$. L'unica relazione possibile è una relazione di proporzionalità la cui costante si determina essere $8\pi G/c^4$ dove G è la costante di gravitazione universale e quindi

$$G_{\mu\nu} = \frac{8\pi G}{c^4} T_{\mu\nu}. \quad (13.14)$$

In assenza di materia ed energia $T_{\mu\nu} = 0$ per ogni coppia μ, ν . In questo caso $G_{\mu\nu} = 0$ il che implica che le variazioni delle $g_{\mu\nu}$ siano nulle per cui $g_{\mu\nu}$ risulta essere costante ed evidentemente deve essere proprio il tensore determinato dalla relatività ristretta.

In presenza di materia, invece, lo spazio-tempo si deforma e la forza di gravità che si esercita tra i corpi non è altro che il risultato di questa deformazione. Come una persona che cammini in linea retta su una superficie curva come la Terra percorre di fatto una traiettoria curva, un corpo che si muove in linea retta in uno spazio curvo come quello prodotto dalla presenza di un'altra massa nell'Universo percorrerebbe una traiettoria che, vista da un ipotetico sistema inerziale, apparirebbe curva.

13.4 Effetti gravitazionali sul tempo

La teoria della relatività ristretta afferma che il tempo non è assoluto, ma scorre in modo diverso nei diversi sistemi di riferimento. La teoria della relatività generale ci dice invece che trovarsi nelle vicinanze di un corpo massivo equivale a trovarsi in un sistema

²L'impulso o quantità di moto dipende dalla massa, dunque dà il contributo della materia a questo tensore

di riferimento non inerziale (quindi sicuramente in moto).

Esercizio 13.1 *Spazi curvi*

Gli orologi atomici che si trovano a bordo dei satelliti GPS sono soggetti a una forza gravitazionale inferiore rispetto a quella cui sono soggetti gli orologi a Terra. Calcola l'anticipo degli orologi di bordo dovuto alla curvatura dello spazio-tempo prodotta dalla massa della Terra.

SOLUZIONE →

Il campo gravitazionale diminuisce come $1/r^2$ con la distanza r dal centro del corpo che lo genera, quindi il campo è più intenso vicino alla superficie della Terra e meno intenso ad alta quota. Un orologio a Terra, dunque, segna il tempo come se si trovasse in un laboratorio più accelerato rispetto a un orologio a bordo di un satellite. Di conseguenza il tempo a Terra scorre più lentamente di quanto non faccia a bordo di un satellite. La teoria della relatività generale permette di calcolare a quanto ammonta il ritardo, ma i calcoli sono molto complicati. Nonostante questo si può ottenere un'ottima stima del ritardo relativo tra gli orologi usando semplici argomenti dimensionali. Se t è il tempo a Terra e τ quello a bordo del satellite, sicuramente deve essere

$$t = \alpha \tau \quad (13.15)$$

dove α è una costante adimensionale. α deve potersi scrivere nella forma

$$\alpha \simeq 1 + \delta \quad (13.16)$$

perché in assenza di campo gravitazionale dev'essere $t = \tau$ e quindi $\delta = 0$. Evidentemente δ deve dipendere dal potenziale del campo gravitazionale³ che sappiamo scrivere come

$$\mathcal{G} = -G \frac{M}{r} \quad (13.17)$$

³Deve dipendere dal potenziale e non dal campo perché deve essere uno scalare, come il potenziale, mentre il campo è un vettore.

dove M è la massa della Terra e r la distanza dal centro. Evidentemente il potenziale ha le dimensioni di un'energia divisa per una massa, quindi ha le dimensioni di una velocità al quadrato. L'unico modo di costruire una grandezza adimensionale quindi consiste nel dividere questa quantità per una velocità al quadrato, che non può che essere l'unica velocità indipendente dai sistemi di riferimento: c . Deve quindi essere

$$t \simeq \left(1 + G \frac{M}{c^2 r}\right) \tau, \quad (13.18)$$

che è praticamente identica al risultato che si trova usando la teoria completa. Il segno $+$ presente in questa relazione indica che il tempo in prossimità di un corpo massivo scorre più lentamente rispetto a quello che scorre in punti dell'Universo lontani da sorgenti di campo gravitazionale. Che il segno sia giusto lo si capisce osservando che in assenza di campi non ci sono accelerazioni ed è come se l'orologio fosse in un sistema di riferimento fermo. In presenza di accelerazioni si misura una forza ed è come se l'orologio si muovesse. Per accelerazioni piccole il sistema si deve comportare come nella relatività ristretta e rispetto al caso dell'orologio fermo deve risultare una durata maggiore dell'intervallo di tempo.

Il paradosso dei gemelli - 2

Il fatto che il tempo scorra in maniera diversa in sistemi di riferimento in moto l'uno rispetto all'altro sembra causare un paradosso perché il risultato di una *misura*, che consiste nel confrontare l'età di uno dei gemelli con quella dell'altro, è diverso secondo il punto di vista adottato. In realtà il paradosso non esiste perché la misura è impossibile dal momento che la relatività ristretta si applica solo nel caso di sistemi di riferimento inerziali e un'astronave che parte e torna successivamente da dov'è partita di sicuro accelera in almeno tre fasi del viaggio (partenza, inversione del moto, arrivo).

La relatività generale risolve il paradosso affermando che in effetti il viaggiatore torna con un'età diversa rispetto a quella del gemello. Mentre nei sistemi di riferimento in moto rettilineo uniforme l'uno rispetto all'altro non è possibile stabilire in maniera oggettiva quale dei due sia in moto e quale sia fermo (ma si può davvero stabilire che un sistema di riferimento è fermo?), nel caso di sistemi di riferimento accelerato si può sempre dire quale dei due sia in moto e quale no: quello nel quale si manifestano le forze apparenti è il sistema in moto. Nel sistema di riferimento del gemello viaggiatore sono presenti accelerazioni e decelerazioni che provocano un'alterazione della marcia degli orologi che dipende dall'intensità dell'accelerazione.

Più è grande quest'accelerazione più lentamente scorre il tempo a bordo dell'astronave. Secondo che la somma delle accelerazioni a bordo sia maggiore o minore di quelle subite dal gemello a terra (che è soggetto alla gravità di quest'ultima), il gemello viaggiatore può tornare più o meno giovane.

Che cos'è lo spazio <http://www.radioscienza.it/2012/09/16/lo-spazio/>

$E = mc^2$ <http://www.radioscienza.it/2013/11/01/emc2/>

Il GPS <http://www.radioscienza.it/2012/06/18/il-gps/>

La teoria della Relatività <http://www.radioscienza.it/2013/01/16/la-teoria-della-relativita/>

La Relatività Generale <http://www.radioscienza.it/2013/06/28/la-relativita-generale/>

Fisicast

I seguenti podcast di **Fisicast** hanno a che fare con l'argomento trattato negli ultimi due capitoli:

La Meccanica Quantistica

PREREQUISITI: cinematica classica e relativistica, momenti magnetici, forze elettriche

Lo studio della meccanica quantistica rappresenta una tappa particolarmente importante per la formazione scientifica. Come nel caso della scoperta dei [raggi cosmici](#), che darà origine alla fisica delle particelle, anche la meccanica quantistica nasce grazie a un evento tutto sommato del tutto secondario che si potrebbe definire irrilevante, se non se ne conoscessero le conseguenze: lo studio dello **spettro di corpo nero**. In fisica, come in altre discipline, non c'è niente di secondario: tutte le domande devono trovare una risposta, qualunque sia il giudizio a priori sull'importanza di queste.

In questa pubblicazione non intendiamo ripercorrere le tappe di questa scoperta da un punto di vista storico. Preferiamo portare il lettore a *scoprire* la meccanica quantistica attraverso una storia leggermente alterata, che comunque include alcune delle scoperte fondamentali che hanno portato alla formulazione di questa teoria.

14.1 Il corpo nero

Il così detto problema del **corpo nero** consiste nel trovare la forma dello **spettro** della radiazione elettromagnetica emessa da un corpo con emissività pari a 1. Detto in questi termini il problema appare difficile e astruso, ma lo si può riformulare in termini meno generali e più familiari in questo modo: perché mai guardando il forno di un pizzaiolo quando è freddo appare nero e quando invece è caldo è di colore giallo?

Predire, in base alle leggi della fisica, il colore del forno conoscendone la temperatura può apparire un problema del tutto irrilevante (si dorme certamente bene anche senza saperlo fare), ma per alcuni si trattava di un problema cui dare una risposta. Si può cercare di trovare una risposta a questa domanda usando le leggi della termodinamica e dell'elettromagnetismo insieme. Riprodurre i calcoli che si facevano al tempo per cercare di trovare la risposta a questa domanda è piuttosto noioso e in fondo non aggiunge nulla alla nostra conoscenza, ma il risultato di tutti i conti è che lo spettro predetto (cioè la distribuzione d'intensità della radiazione in funzione della sua lunghezza d'onda) non è in accordo con i dati sperimentali.

Cerchiamo di capire perché. Intanto osserviamo che non tutti i corpi che emettono luce sono caldi (un LED, ad esempio), ma tutti i corpi caldi emettono luce (o comunque una qualche forma di radiazione elettromagnetica). Un pezzo di carbone scaldato emette una viva luce rossa e i corpi umani, che si trovano a una temperatura di 36 °C emettono radiazione infrarossa, invisibile per i nostri occhi, ma visibile attraverso opportuni dispositivi sensibili a quel tipo di radiazione.

Consideriamo una singola particella elettricamente carica nel vuoto, che venga investita da una radiazione elettromagnetica. Il campo elettrico che costituisce questa radiazione fa muovere la particella, che oscilla in sincrono con il campo. Una particella carica accelerata, però, emette radiazione elettromagnetica: funziona infatti come un'antenna. La frequenza della radiazione emessa è quella di oscillazione che a sua volta è quella della radiazione incidente.

Le particelle che formano un gas, un liquido o un solido, però non sono libere di muoversi. In ge-

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.1 I visori notturni funzionano amplificando e trasformando in luce visibile la radiazione infrarossa emessa con maggiore intensità dai corpi più caldi, che si comportano, in buona approssimazione, come corpi neri. Il filmato è una gentile concessione di FLIR® [<https://www.youtube.com/watch?v=rAvnMYqj2c0>].

nerale le particelle sono legate le une alle altre da forze che possiamo comunque pensare, in prima approssimazione, funzionare come molle che collegano le varie particelle. Nel caso del gas più semplice, l'idrogeno, le particelle coinvolte sono due: un protone positivo e un elettrone negativo, legate tra loro da forze elettrostatiche. Se una radiazione elettromagnetica investe un atomo d'idrogeno fa muovere gli elettroni e i protoni, che tuttavia non possono allontanarsi più di tanto l'uno dall'altro. L'ampiezza dell'oscillazione può essere più o meno grande, secondo le caratteristiche della forza che lega le particelle. Per questo i corpi talvolta non riescono ad assorbire l'energia della radiazione incidente e altre volte, pur assorbendola, non la riemettono con le stesse modalità. Un **corpo nero** è un corpo che assorbe integralmente tutta l'energia che riceve, senza rimetterla.

Le particelle che formano un corpo si possono mettere in moto anche senza l'ausilio di un'onda elettromagnetica. Se riscaldiamo qualcosa come un pezzo di carbone, le sue particelle si mettono a oscillare sempre più rapidamente, perché la loro energia cinetica media aumenta. Anche le particelle del carbone in fin dei conti sono cariche elettricamente e quindi oscillando attorno a una posizione d'equilibrio irradiano radiazione elettromagnetica. All'inizio la frequenza di oscillazione è bassa, e quindi il corpo, anche se illuminato, appare nero, perché la

radiazione incidente non è sufficiente a farlo irraggiare. Ma all'aumentare della temperatura aumenta la frequenza di oscillazione e le particelle cominciano a irradiare una radiazione a frequenze sempre più alte a al limite visibili dall'occhio umano: il carbone diventa rosso.

Quello che può succedere è che la radiazione prodotta dal carbone fa muovere altre particelle dello stesso carbone che potrebbero già essere in oscillazione. Se le oscillazioni non sono in fase il risultato è un movimento caotico delle particelle con varie frequenze. La somma delle diverse frequenze dà origine al colore che osserviamo. Il colore è determinato da una condizione di equilibrio tra la radiazione emessa e quella assorbita dalle particelle in moto. Tutte le frequenze sono permesse e l'ampiezza dell'oscillazione può variare da zero a un valore massimo.

In linea di principio si possono eseguire dei calcoli per predire la forma di questo spettro: per un corpo non nero il calcolo è complicato dal fatto che non tutta la radiazione emessa è dovuta a questo meccanismo. Un corpo che a temperatura ambiente appare verde possiede la proprietà di diffondere la luce verde che vi incide a causa delle proprietà ottiche della sua superficie. Per un corpo nero invece questi effetti si possono trascurare. I calcoli esatti portano a un risultato incomprensibile: si ottiene sempre che l'energia irraggiata da un corpo nero caldo è di fatto infinita (come si vede nella Figura 14.2 al diminuire della lunghezza d'onda la quantità di radiazione emessa tende rapidamente a infinito). Questo è impossibile!

Nel 1901 **Max Planck** pubblica un articolo [12] nel quale esegue i conti facendo l'ipotesi semplificativa secondo la quale la radiazione può essere emessa solo in quantità discrete chiamate **quanti**. L'idea è la seguente: per semplificare i conti assumo che le particelle cariche nel carbone possano oscillare a una sola frequenza. Oscillando producono radiazione elettromagnetica a quell'unica frequenza. Questa radiazione potrebbe essere assorbita da una delle particelle del pezzo di carbone oppure no, quindi, se sono stati prodotti N quanti dal carbone, da questo possono uscirne $0, 1, 2, \dots, n$ con $n \leq N$, mentre i restanti

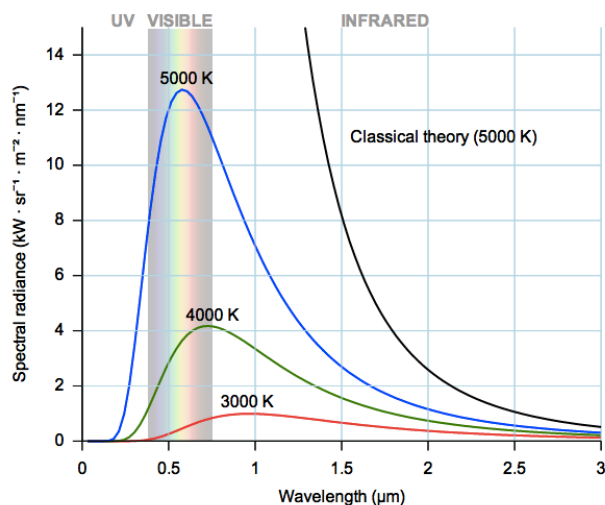


Figura 14.2 Lo spettro di corpo nero per diverse temperature del corpo. Nella figura si vede anche la previsione della teoria classica della luce. In ascissa c'è la lunghezza d'onda della luce emessa, mentre in ordinata la potenza emessa per unità di angolo solido per metro quadro.

$N - n$ sono assorbiti e contribuiscono al moto delle particelle di cui è composto il carbone. In questo modo la somma delle energie trasportate dai singoli quanti è certamente finita (non può essere maggiore di NE dove E è l'energia di un singolo quanto). E dal momento che l'energia irradiata dipende dalla frequenza della radiazione emessa E deve essere proporzionale alla frequenza ν : $E = h\nu$ dove h è una costante (che poi prenderà il nome di **costante di Planck**) che vale 6.63×10^{-34} J s o 4.14×10^{-15} eV s.

Questo modello spiega bene anche un'altra questione: se illumino un corpo (nero, ma non solo) con la luce, il corpo si scalda. Se si scalda vuol dire che assorbe energia e che le sue particelle si mettono in moto. Dovrebbe quindi irraggiare e in effetti è così, ma dovrebbe farlo a una frequenza del tutto identica a quella della luce che lo illumina. Di solito irraggia a una frequenza più bassa (tipicamente nell'infrarosso). Se assumiamo che la luce sia fatta di corpuscoli detti quanti, questi possono colpire l'elet-

trone di un atomo in un determinato istante *allontanandolo* dal nucleo. Una volta allontanato, prima di essere nuovamente colpito da un quanto, l'elettrone è richiamato dalla forza elastica che lo trattiene al nucleo e dunque si muove eseguendo oscillazioni a frequenza più bassa rispetto a quelle tipiche della radiazione incidente.

La teoria di Planck non ebbe molto successo: risolveva sì alcuni problemi, ma usando un *trucco*. Non si poteva considerare una soluzione.

14.2 L'effetto fotoelettrico

Nel 1887, prima dell'ipotesi di Planck, un fisico di nome **Heinrich Hertz** stava conducendo alcuni esperimenti per studiare la propagazione delle onde elettromagnetiche. Per farlo usava uno strumento che produceva scintille quando arrivavano onde elettromagnetiche. Per vederle meglio, Hertz cercò di fare il maggior buio possibile nel suo laboratorio e si accorse che la qualità delle scintille peggiorava: si vedevano meglio, ma erano più deboli. Le scintille sono provocate dal rapido movimento delle cariche elettriche (come sopra, le cariche elettriche accelerate emettono radiazione elettromagnetica, quindi luce se l'emissione avviene a frequenze opportune), quindi Hertz ne concluse giustamente che nei suoi strumenti le cariche si muovevano più facilmente se illuminati. In particolare notò che le scintille erano molto intense in presenza di radiazione ultravioletta, un po' meno intense se s'illuminava tutto con luce blu e praticamente assenti se gli strumenti venivano illuminati da luce di colore giallo o rosso. Solo i corpi carichi negativamente subivano quest'effetto, poi chiamato **effetto fotoelettrico**. Inoltre, quando l'illuminazione provocava la scarica, questa era tanto più intensa quanto maggiore era l'intensità della luce.

L'effetto fotoelettrico, in sostanza, ha le seguenti caratteristiche:

- consiste nella perdita di carica elettrica da parte di un corpo carico, quando questo è illuminato;

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.3 L'effetto fotoelettrico si spiega facilmente se si ammette che la luce sia composta di corpuscoli invece che di onde. <https://www.youtube.com/watch?v=ZvyFusOnC6s>

- avviene solo se il corpo è carico negativamente;
- avviene solo se la frequenza della radiazione che lo provoca è superiore a una determinata soglia ν_0 che dipende dal materiale di cui è fatto il corpo;
- se avviene, l'intensità della scarica è proporzionale all'intensità della radiazione.

Un modo per spiegarlo è il seguente: un corpo elettricamente carico di carica negativa possiede un eccesso di elettroni, debolmente legati al corpo. Quando una radiazione elettromagnetica raggiunge gli elettroni, trasferisce a essi una certa quantità d'energia. Quest'energia potrebbe essere sufficiente a liberare gli elettroni dai deboli legami che li trattengono sul corpo e, naturalmente, quanto maggiore è l'intensità della radiazione incidente, tanto più alto sarà il numero di elettroni che riescono a guadagnare l'energia necessaria e ad abbandonare il corpo, scaricandolo. Il ragionamento sembra non fare una piega: in fondo le onde trasportano energia! Ma l'energia trasportata da un'onda dipende dalla sua ampiezza, non dalla sua frequenza! Non si spiega, dunque, come mai l'effetto fotoelettrico si manifesti solo per onde di frequenza maggiore di una frequenza caratteristica del materiale e non per onde di frequenza più bassa, ma di ampiezza grande a piacere. In altre parole una luce gialla molto intensa dovrebbe trasportare molta più energia di una luce blu poco intensa e invece l'effetto fotoelettrico si manifesta con la seconda, ma non con la prima.

Esercizio 14.1 *Effetto fotoelettrico*

La maggior parte della luce che proviene dal Sole ha una lunghezza d'onda compresa tra i 400 nm e i 750 nm. Calcola quanta energia possono al massimo guadagnare gli elettroni del materiale illuminato per effetto fotoelettrico.

SOLUZIONE →

L'unica soluzione possibile è quella di ammettere che la luce non sia un'onda, ma un flusso di corpuscoli di energia E proporzionale alla frequenza associata all'onda elettromagnetica corrispondente: $E = h\nu$. Fu **Albert Einstein** a proporre questa soluzione [13], in uno dei suoi tre articoli più importanti (tutti e tre pubblicati nell'*annus mirabilis* 1905). Per questa scoperta, tra l'altro, ad Einstein fu conferito il Premio Nobel (con sua grande disapprovazione, dal momento che riteneva di meritarlo per la teoria della relatività e non per quel che considerava tutto sommato un errore: la sua *Nobel Lecture*, tradizionalmente riservata all'argomento oggetto del premio, fu infatti sulla relatività e non sull'effetto fotoelettrico). A questi corpuscoli o quanti di luce si dà il nome di **fotoni**.

Si viene così a determinare un apparente paradosso: gli esperimenti mostrano in maniera inequivocabile che la luce è un'onda (attraversando una fenditura produce il fenomeno della **diffrazione**, tipico delle onde). Ma lo spettro di corpo nero e l'effetto fotoelettrico si spiegano solo ammettendo che sia composta di corpuscoli di energia $E = h\nu$. Come si risolve il *controsenso*? È semplice: ammettendo che la luce sia al contempo un'onda e una particella. La soluzione può apparire artificiosa e quasi un *trucco*, essendo completamente priva di logica. Ma chi dice che l'Universo segua la *nostra* logica? Come nel caso della relatività, l'Universo funziona secondo sue proprie Leggi per le quali non esiste nessun motivo di principio secondo il quale queste Leggi debbano uniformarsi al nostro modo (alla nostra capacità) di pensare le cose. Il fatto che a noi sembri assurdo che qualcosa possa essere al contempo un'onda

e una particella non è un buon motivo per rifiutare questa spiegazione.

Si potrebbe anche pensare che l'apparente paradosso derivi da una nostra scarsa conoscenza dei fenomeni per cui, in effetti, la luce dev'essere una delle due cose, ma per motivi a noi ignoti si comporta *come se* fosse l'altra in certe occasioni. Per la Fisica, però, *comportarsi come* o *essere* sono sinonimi. Inoltre, numerosi altri fenomeni, riportati più sotto, evidenziano che in effetti la spiegazione che ci siamo dati dev'essere quella giusta: la luce non è un'onda o una particella: è tutt'e due le cose allo stesso tempo! Per quanto strano possa sembrare non possiamo che ammettere che dev'essere così: sono i risultati degli esperimenti che stabiliscono la natura delle cose e per la luce questo è il risultato. Citando ancora **Luigi Malerba** "Succede purtroppo che spesso i fatti smentiscono le ingegnose e confortevoli teorie mentre non si sono mai viste teorie che smentiscono i fatti"¹.

Accettato questo punto di vista, la luce di frequenza ν si deve considerare come un flusso di fotoni di energia $E = h\nu$. Sapendo che $\omega = 2\pi\nu$ possiamo anche scrivere $E = h\frac{\omega}{2\pi} = \hbar\omega$ avendo definito $\hbar = h/2\pi$. La lunghezza d'onda λ è legata alla frequenza dalla relazione $\lambda = c/\nu$, perciò possiamo anche scrivere che $E = hc/\lambda$. La quantità di moto trasportata da un fotone, inoltre, è in unità naturali uguale alla sua energia $E = h\nu$, ma in unità SI, essendo $E^2 = p^2c^2 + m^2c^4$ e $m = 0$, è $p = E/c = h/\lambda$. Ricordando, infine, che $k = 2\pi/\lambda$, $p = \hbar k$.

14.3 L'effetto Compton

L'effetto Compton consiste nella diffusione anelastica della luce da parte dei materiali. Se si invia un fascio luminoso su un materiale, la luce diffusa da questo ha la stessa lunghezza d'onda di quella incidente, per gli stessi motivi addotti sopra per spiegare il meccanismo di emissione di un corpo nero.

In alcuni casi però la luce diffusa si presenta con una lunghezza d'onda maggiore e la differenza rispetto a quella incidente dipende dall'angolo di dif-

fusione. Questo fenomeno si spiega molto bene con l'ipotesi quantistica. La luce infatti si può considerare come un flusso di fotoni che *piove* addosso agli elettroni degli atomi di cui è costituito il materiale. I fotoni che urtano questi elettroni possono essere diffusi ad angoli diversi senza perdere la loro energia, ma potrebbe anche accadere che parte dell'energia dei fotoni è ceduta agli elettroni del materiale. Se l'energia dei fotoni è abbastanza alta, gli elettroni si possono considerare come liberi (non legati ai nuclei): questa approssimazione si traduce nella condizione

$$h\nu \gg V \quad (14.1)$$

dove V rappresenta l'energia di legame degli elettroni negli atomi. In questo caso possiamo trattare gli elettroni come particelle libere e i fotoni come particelle che possono urtarli e cedere loro parte dell'energia iniziale. Usando la **cinematica relativistica** il conto è semplicissimo: in **unità naturali** il quadrimpulso dei fotoni incidenti è $\underline{P}_\gamma = (E, \mathbf{p})$, con $E = p$ (indicando con p il modulo del vettore $\mathbf{p}$), e quello degli elettroni nello stato iniziale è $\underline{P}_e = (m, \mathbf{0})$ dal momento che si possono considerare fermi (se non lo sono basta mettersi nel sistema di riferimento in cui lo sono). Dopo l'urto il quadrimpulso dei fotoni cambia e diventa $\underline{P}'_\gamma = (E', \mathbf{p}')$, così come quello degli elettroni che possiamo indicare con $\underline{P}'_e = (E_e, \mathbf{p}_e)$. Il quadrimpulso totale si conserva perciò

$$\underline{P}_\gamma + \underline{P}_e = \underline{P}'_\gamma + \underline{P}'_e. \quad (14.2)$$

Dal momento che siamo interessati a conoscere la cinematica del fotone nello stato finale, mentre di quella dell'elettrone non c'interessa nulla possiamo isolare il quadrimpulso dell'elettrone

$$\underline{P}_\gamma + \underline{P}_e - \underline{P}'_\gamma = \underline{P}'_e. \quad (14.3)$$

e fare il modulo quadro di entrambi i membri per ottenere

$$P_\gamma^2 + P_e^2 + P_\gamma'^2 + 2\underline{P}_\gamma \cdot \underline{P}_e - 2\underline{P}_\gamma \cdot \underline{P}'_\gamma - 2\underline{P}_e \cdot \underline{P}'_\gamma = P_e'^2. \quad (14.4)$$

¹Luigi Malerba, La superficie di Eliane, ed. Mondadori

Ricordando che il modulo quadro del quadrimpulso è un invariante relativistico pari alla massa a riposo della particella al quadrato abbiamo che $P_\gamma^2 = P_\gamma'^2 = 0$ e $P_e^2 = P_e'^2 = m^2$, perciò sostituendo

$$2\underline{P}_\gamma \cdot \underline{P}_e - 2\underline{P}_\gamma \cdot \underline{P}_\gamma' - 2\underline{P}_e \cdot \underline{P}_\gamma' = 0. \quad (14.5)$$

Il prodotto scalare dei quadrivettori si esegue moltiplicando le coordinate omologhe con il segno $-$ se sono spaziali e $+$ se sono temporali; dividendo tutto per 2 e svolgendo i prodotti si trova che

$$Em - (EE' - pp' \cos \theta + mE') = 0 \quad (14.6)$$

dove θ è l'angolo compreso tra $\mathbf{p}$ e $\mathbf{p}'$, cioè l'angolo di diffusione. Ricordando che $E = p$ nonché $E' = p'$ si trova che

$$m(E - E') - EE'(1 - \cos \theta) = 0 \quad (14.7)$$

cioè che, raccogliendo E' e portando E all'altro membro,

$$E' = \frac{mE}{E(1 - \cos \theta) + m} = \frac{m}{(1 - \cos \theta) + \frac{m}{E}}. \quad (14.8)$$

È interessante studiare i limiti di quest'espressione. Per $\theta = 0$, $E' = E$ come ci si aspetta non essendoci alcuna diffusione. Quando invece $\cos \theta = -1$, cioè per fotoni diffusi all'indietro, si trova che

$$E' = \frac{m}{2 + \frac{m}{E}} = \frac{mE}{2E + m} \quad (14.9)$$

dalla quale si evince che E' non può mai essere nulla (a meno che $E = 0$, che è l'energia di un fotone di lunghezza d'onda infinita che non interagisce), cioè che il fotone non può mai perdere tutta la sua energia per effetto Compton.

Se dividiamo l'equazione (14.7) per EE' possiamo riscrivere la relazione che lega E' a E e all'angolo come

$$\frac{1}{E'} - \frac{1}{E} = \frac{1 - \cos \theta}{m} \quad (14.10)$$

che è comoda perché permette di riscriverla in termini di frequenza e lunghezza d'onda. Infatti, ricordando che $E = h\nu$ e che $E = h\nu'$ abbiamo

$$\frac{1}{\nu'} - \frac{1}{\nu} = h \frac{1 - \cos \theta}{m} \quad (14.11)$$

e scrivendo $\nu = c/\lambda$, $\nu' = c/\lambda'$

$$\lambda' - \lambda = h \frac{1 - \cos \theta}{mc} = \lambda_C (1 - \cos \theta) \quad (14.12)$$

dove $\lambda_C = h/mc$ è chiamata **lunghezza d'onda Compton dell'elettrone** e ha, come deve, le dimensioni fisiche di una lunghezza.

14.4 La misura e il Principio d'indeterminazione

Con la meccanica relativistica si era già capito che la fisica non poteva aver a che fare con quantità assolute, ma con i risultati di una misura, che sono relativi all'osservatore. L'avvento della meccanica quantistica rafforza questa convinzione: non possiamo considerare qualcosa come avente una natura di un tipo o dell'altro in maniera assoluta. Dipende da come si esegue la misura.

L'esecuzione di una misura consiste sempre nel far interagire l'oggetto da osservare con uno strumento. Il fatto che ci sia un'interazione implica l'esistenza di forze che si producono tra strumento e oggetto osservato, per cui lo stato dell'oggetto osservato può risultare alterato dall'applicazione di queste forze. La conseguenza è che quello che si osserva non è mai lo stato in quanto tale, ma lo stato che l'oggetto possiede avendo interagito con lo strumento.

In meccanica classica lo stato di un punto materiale è perfettamente definito quando se ne conoscano contemporaneamente e con precisione arbitraria la posizione e la velocità. La **fisica relativistica** insegna che la velocità v di una particella di massa m non è una buona misura dello stato in quanto in seguito all'applicazione di una forza quel che cambia non è la velocità della particella, ma la sua quantità di moto $p = mv$ (per velocità $v \ll c$, p è la stessa

cosa di v a parte un fattore costante, ma in meccanica relativistica $m = m(v)$). La natura del principio classico però non cambia: per conoscere lo stato di una particella basta conoscerne con precisione infinita la posizione x e la quantità di moto p . Che tecnicamente non sia fattibile (non esistono strumenti che possano misurare una grandezza fisica con precisione arbitraria) non importa: quel che importa è che in linea di principio si possa fare: se conoscessi **esattamente** la posizione e la quantità di moto di un punto materiale al tempo t potrei conoscerne lo stato a un tempo t' successivo (o precedente) a t .

Proviamo ad applicare questo principio all'osservazione di una particella molto piccola come un elettrone legato in un atomo. Per poterne misurare la posizione si potrebbe **illuminare** l'elettrone con una luce opportuna: osservando la luce diffusa da questo se ne ricaverebbe la posizione (che poi è il metodo con il quale osserviamo qualunque cosa con i nostri occhi: al buio non si vede nulla; è la luce diffusa dagli oggetti che ci dice dove sono). Se però la lunghezza d'onda della luce non è abbastanza piccola, il moto dell'onda elettromagnetica non viene perturbato affatto dalla presenza dell'elettrone che è come se non ci fosse. Per osservarlo è necessario usare una lunghezza d'onda sufficientemente piccola: usare quindi una radiazione elettromagnetica come i raggi X o i raggi γ . Ma in questi casi la luce diventa un flusso di particelle di energia $E = h\nu$ che interagiscono con l'elettrone urtandolo e cambiandone quindi la velocità (o meglio la quantità di moto) in seguito all'urto. La misura della posizione dell'elettrone implica l'alterazione del suo stato di moto, dunque è impossibile conoscerne posizione e quantità di moto contemporaneamente. Una delle due grandezze fisiche non è determinabile al di là di una data precisione e non in virtù di un limite tecnologico, ma per ragioni intrinseche al processo di misura.

Il Principio d'indeterminazione afferma proprio questo: il prodotto tra l'indeterminazione sulla posizione e quella sulla quantità di moto di qualunque oggetto non può mai essere nullo! Deve essere sempre maggiore di un numero che, per quanto piccolo, non è mai zero. Se, ad esempio, illuminiamo un

elettrone con una luce di frequenza ν e con numero d'onda k , trasferiamo all'elettrone una quantità di moto (vedi a pag. 157) $\Delta p = \hbar k = \frac{2\pi}{\lambda}$. La lunghezza d'onda λ dev'essere $\lambda \lesssim \Delta x$ per localizzare l'elettrone dunque

$$\lambda = \frac{h}{\Delta p} \lesssim \Delta x \quad (14.13)$$

cioè che

$$\Delta x \Delta p \gtrsim h. \quad (14.14)$$

Il principio² si può riscrivere in un altro modo, dividendo tutto per Δt :

$$\frac{\Delta x}{\Delta t} \Delta p = v \Delta p = 2\Delta E \geq \frac{h}{\Delta t} \quad (14.15)$$

e quindi

$$\Delta E \Delta t \geq \frac{h}{2}. \quad (14.16)$$

Questa relazione indica che misurando l'energia di un particolare stato, questa può avere un'indeterminazione ΔE che dipende dal tempo trascorso per eseguirla. Consideriamo un caso particolare, che consiste nella determinazione della massa di una particella che, come sappiamo, è una misura dell'energia a riposo dello stato. Se la misura si può protrarre per molto tempo Δt tende a infinito, quindi ΔE tende a zero e la misura si può eseguire con precisione arbitraria. È il caso, ad esempio, della misura della massa di un elettrone o di un protone, che assume un valore perfettamente determinato anche se, a causa delle limitazioni tecnologiche, avrà comunque un'indeterminazione di natura statistica.

Se invece misuriamo la massa di una particella instabile come un **muone**, la misura non si può eseguire a un tempo arbitrario: deve essere eseguita entro il tempo di vita della particella. Trascorso tale tempo la particella non esiste più e non se ne può di

²In letteratura si trovano espressioni che possono differire di piccoli fattori da questa, che è stata ricavata attraverso un'approssimazione un po' grossolana; di solito si trova $\Delta x \Delta p \geq \hbar$ o $\Delta x \Delta p \geq \hbar/2$. Quel che conta è però l'ordine di grandezza, insieme alle conseguenze dell'indeterminazione.

certo misurare la massa. In questo caso la sua energia a riposo possiede un'indeterminazione **intrinseca** che ne fa fluttuare la misura di massa entro un limite dell'ordine di $h/2\tau$, dove τ è il tempo di decadimento. Non si tratta di un effetto statistico, ma di un fenomeno quantistico intrinseco nella definizione di massa della particella. Se disponessimo di uno strumento preciso al di là del limite quantistico troveremmo che non tutti i muoni hanno la stessa massa, ma che ognuno ne avrebbe una diversa con una certa distribuzione (simile, anche se non proprio identica, a una gaussiana), la cui larghezza non è determinata da effetti di natura stocastica, ma dalle fluttuazioni quantistiche che determinano l'istante nel quale la particella decade. L'indeterminazione intrinseca della massa di una particella si chiama talvolta **larghezza** della particella: una particella **stretta** è una particella con una vita media lunga (ΔE piccolo, Δt grande), mentre una particella **larga** ha una vita media molto breve.

14.5 Onde di materia

La luce è chiaramente un'onda in certi contesti, ma diventa una particella in altre condizioni. Perché dunque non potrebbe avvenire il contrario? Che quelle che consideriamo naturalmente particelle possano comportarsi come onde in certi contesti? In effetti questo è quel che accade. La frequenza ν di un fotone è legata alla sua quantità di moto p secondo la legge $p = h\nu$. Una particella di quantità di moto p dunque potrebbe comportarsi come un'onda di frequenza

$$\nu = \frac{p}{h}. \quad (14.17)$$

Essendo $\nu = c/\lambda$ dove c è la velocità dell'onda (che in questo caso coincide con la velocità della luce) e λ la sua lunghezza d'onda, possiamo anche scrivere che

$$\frac{\lambda}{c} = \frac{h}{p} \quad (14.18)$$

che in **unità naturali** diventa $\lambda = h/p$ (e ponendo anche $h = 1$, $\lambda = 1/p$).

Se fosse vero si potrebbero realizzare esperimenti nei quali un fascio di particelle (per esempio di elettroni), attraversando un reticolo dovrebbe produrre al di là di esso una figura d'interferenza! Questo effettivamente è ciò che accade. Se si invia un fascio di elettroni prodotti da un tubo catodico su un reticolo molto fine, sullo schermo del tubo catodico non si vede l'immagine delle fenditure come ci si aspetterebbe se gli elettroni fossero particelle, ma una figura di diffrazione con massimi e minimi: segno che gli elettroni si comportano come onde di una certa lunghezza d'onda che interferiscono tra loro attraversando le fenditure del reticolo. Si tratta di una tecnica ormai consolidata che permette, per esempio, di studiare la forma dei reticoli cristallini dei materiali facendo attraversare un campione di quel materiale da un fascio di particelle di energia opportuna, tale da produrre **onde di materia** di lunghezza comparabile con quella del passo reticolare.

Esercizio 14.2 *Microscopi elettronici*

In un microscopio elettronico gli elettroni sono accelerati da una differenza di potenziale di 80 kV. Quanto è migliore il potere risolutivo del microscopio elettronico rispetto a quello ottico?

SOLUZIONE →

Grazie a questo *dualismo onda-particella* si possono realizzare strumenti come i **microscopi elettronici** in cui fasci di elettroni attraversano i campioni da analizzare. La lunghezza d'onda di questi fasci è molto minore di quella della luce e questo permette di ottenere risoluzioni molto più spinte rispetto ai microscopi ottici.

14.6 Gli atomi

Secondo la fisica classica un atomo potrebbe funzionare come un sistema solare in miniatura. Tuttavia, per la meccanica quantistica non è così: gli atomi infatti non si possono osservare con luce visibile per-

ché hanno dimensioni troppo piccole. Per **osservarli** occorre farli interagire con una radiazione di frequenza più alta che però avrà energia più alta e produrrà una perturbazione rilevante sul sistema. Non si può osservare un atomo, ma se ne può misurare lo stato quanto interagisce con la radiazione. Misurando la radiazione prodotta dall'interazione se ne può determinare l'energia, ad esempio. Conoscendo l'energia della radiazione incidente possiamo quindi determinare l'energia assorbita dall'atomo (o meglio da uno o più dei suoi elettroni).

Non ha alcun senso chiedersi *dove* sia un elettrone rispetto al nucleo in un dato sistema di riferimento perché non si può misurare questa quantità. Si può misurare l'energia e quindi ci si può chiedere quale sia l'energia dell'elettrone. Lo stato di un elettrone dunque non è caratterizzato dalla sua posizione e dalla sua velocità, ma dall'energia (e forse da qualcos'altro).

Se un atomo fosse fatto come un sistema solare nulla vieterebbe a due o più elettroni di stare sulla stessa orbita, e niente impedirebbe a quest'orbita di avere un raggio qualunque. Al contrario sarebbe vietato a due elettroni risiedere esattamente nello stesso punto nello stesso istante (se c'è uno non c'è posto per l'altro): **due elettroni non possono stare contemporaneamente nello stesso stato**.

14.6.1 Gli spettri atomici

Seguendo le prescrizioni di De Broglie possiamo ammettere che un elettrone non sia assimilabile a un punto materiale, ma a un'onda. L'elettrone, in sostanza, non può trovarsi in un determinato punto dello spazio, ma deve in qualche modo *circondare* il nucleo completamente. Non dovremmo immaginare un elettrone come un punto che si muove attorno al nucleo, ma piuttosto come una specie di *oceano* che ricopra completamente la superficie del nucleo, con le sue onde in superficie. In un atomo ci possono essere più d'uno di questi **oceani elettronici** che si possono anche compenetrare l'uno con l'altro. È come se ci fosse un oceano sopra un altro oceano, la cui superficie attraversa quella dell'altro.

L'onda costituita dalla superficie di questo oceano deve essere stazionaria, perché la superficie deve essere continua. Non può essere che in un punto la superficie di quest'oceano si trova a una certa quota e nel punto adiacente sia molto più alta o molto più bassa. L'onda dunque deve avere una lunghezza d'onda λ discreta:

$$n\lambda = n\frac{h}{p} = 2\pi r \quad (14.19)$$

da cui si ricava che

$$n\frac{h}{2\pi} = pr = mvr. \quad (14.20)$$

Chiamando $\hbar = h/2\pi$ si ha che $mvr = n\hbar$. Il prodotto mvr non è altro che il momento angolare classico dell'elettrone, che dunque risulta **quantizzato**: può solo essere un multiplo di $\hbar$. Se il momento angolare è quantizzato lo è anche l'energia. Prendendo un atomo d'idrogeno (il più semplice di tutti) l'energia potenziale di un elettrone nel campo del nucleo è

$$U = -k\frac{e^2}{r} \quad (14.21)$$

e quella cinetica vale $K = \frac{1}{2}mv^2$. Se il moto è circolare uniforme abbiamo che

$$m\frac{v^2}{r} = k\frac{e^2}{r^2} \quad (14.22)$$

per cui

$$mv^2 = k\frac{e^2}{r} \quad (14.23)$$

e di conseguenza $K = \frac{1}{2}mv^2 = -U/2$. In definitiva $E = K + U = -U/2 + U = U/2$ e quindi

$$E = -k\frac{e^2}{2r}. \quad (14.24)$$

Ora abbiamo che $r = n\hbar/mv$ e che $v = \sqrt{-2E/m}$ e sostituendo:

$$E = -k\frac{e^2}{2n\hbar}m\sqrt{\frac{-2E}{m}}. \quad (14.25)$$

Eleviamo al quadrato e otteniamo



Figura 14.4 Spettro della luce bianca che ha attraversato una sostanza, la quale ne ha assorbito certe frequenze che appaiono come bande scure nello spettro.

$$E^2 = -k^2 \frac{e^4}{2n^2 \hbar^2} m E \quad (14.26)$$

che divisa per E dà

$$E = -k^2 \frac{e^4}{2n^2 \hbar^2} m. \quad (14.27)$$

Le energie di un elettrone in un atomo non possono essere qualunque. Sono permesse solo le energie pari a una costante divisa per n^2 . Si dice che l'energia di un elettrone in un atomo è **quantizzata**.

Questa circostanza spiega in modo naturale l'osservazione sperimentale secondo la quale gli spettri di assorbimento e di emissione della luce si presentano con delle **righe**. Se si invia luce bianca su un recipiente trasparente che contiene un gas di una particolare specie e, dopo che la luce abbia attraversato il materiale si scompone nei suoi colori attraverso un prisma di vetro o un reticolo, si osserva la *manca* di certi colori (Figura 14.4). Lo **spettro** di assorbimento presenta righe scure in corrispondenza di certi valori della lunghezza d'onda della luce caratteristici della specie atomica contenuta nel recipiente. È così, tra l'altro, che possiamo sapere di cosa siano composti oggetti lontani come le stelle: osservandone lo spettro della luce si vedono righe scure in corrispondenza di certe frequenze che sono di fatto la *firma* degli atomi di cui le stelle sono composte.

Le **righe spettrali** si presentano sempre in modo tale da essere distanziate l'una dall'altra di una quantità che diminuisce come $1/n^2$. Alla luce di quanto sopra è facile spiegare perché: la luce giunge sugli atomi in quanti di energia $E = h\nu$. Il singolo

quanto può essere assorbito o meno: è impossibile sottrargli energia parzialmente perché il quanto è indivisibile. Se questa energia coincide con la differenza di energia tra due possibili stati dell'elettrone dell'atomo il quanto può essere assorbito: in conseguenza di ciò l'energia del fotone è ceduta interamente all'elettrone che passa così da uno stato energetico a uno stato energetico diverso, più elevato. I fotoni di energia pari alla differenza di energia tra due livelli atomici dunque non emergono dal recipiente e in corrispondenza di quelle frequenze osserviamo delle righe scure (assenza di luce di quel colore). Se l'energia è più bassa o più alta rispetto alle differenze di cui sopra, il fotone non può essere assorbito perché l'elettrone non può assumere un'energia pari a quella che aveva inizialmente più quella del fotone. È perciò costretto a restare nel suo stato. Di conseguenza il fotone attraversa indenne il recipiente e noi lo osserviamo provocare sullo spettro una riga di colore pari a quello che compete alla frequenza corrispondente.

14.7 Quantizzazione del momento angolare

Il momento angolare L di un elettrone in un atomo è qualcosa che, almeno in linea di principio, si può misurare facendo interagire l'elettrone con la luce. I risultati che possiamo avere da questa misura sono sempre del tipo $L = n\hbar$. In altre parole i possibili valori di L sono sempre multipli interi di $\hbar$. Il momento angolare è un vettore e come tale può avere un'orientazione qualsiasi nello spazio. Scelta una direzione potremmo misurare la proiezione del momento angolare su quella direzione. Ma qualunque misura facciamo non possiamo che ottenere un multiplo di $\hbar$! Altri valori non sono ammessi per il momento angolare. In ogni caso il valore massimo che potremo misurare è L , quindi nel caso $L = 1$ la proiezione di questo vettore su una qualunque direzione può assumere soltanto i valori $m = +1$, $m = 0$ oppure $m = -1$. Questo accade qualunque sia la scelta della direzione rispetto alla quale decidiamo di misurare il momento angolare.

Quando $L = 2$ i possibili valori che si possono misurare per quella che si chiama la **terza componente del momento angolare**³ sono $m = +2$, $m = +1$, $m = 0$, $m = -1$, $m = -2$. In generale se il momento angolare di un elettrone in un atomo è L si possono misurare valori di proiezione del momento angolare lungo una qualunque direzione compresi tra $-L$ e L e separati di ± 1 .

Complessivamente, per un elettrone nello stato di momento angolare L esistono $2L + 1$ possibili stati diversi della terza componente.

14.8 Lo spin degli elettroni

Per quanto un elettrone non sia affatto assimilabile a una carica puntiforme che ruota attorno a un nucleo positivo, possiede pur sempre una carica elettrica e di certo non è qualcosa di statico, di fermo (se fosse assolutamente fermo avrebbe $p = \Delta p = 0$ e la sua posizione sarebbe completamente indeterminata per il Principio d'indeterminazione). Ci si aspetta dunque che possa risentire dell'effetto dei campi magnetici. Un elettrone *classico* è assimilabile a una corrente che circola lungo una spira che ha un **momento magnetico** proporzionale al suo momento angolare, quindi anche per gli elettroni quantistici dovremmo aspettarci qualcosa del genere. La differenza dovrebbe consistere in una quantizzazione del momento magnetico.

Cerchiamo di capire come si muove un atomo con momento magnetico $\boldsymbol{\mu} = (0, 0, \mu_z)$ in un campo magnetico non uniforme. La sua variazione di energia ΔU è $\Delta U = \boldsymbol{\mu} \cdot \Delta \mathbf{B} = \mu_z \Delta B_z$. Quest'energia è acquistata a spese del lavoro fatto dalle forze cui l'atomo è soggetto nel campo magnetico $L = \mathbf{F} \cdot \Delta \mathbf{s}$. Abbiamo dunque che

$$\mu_z \Delta B_z = \mathbf{F} \cdot \Delta \mathbf{s} = F_x \Delta x + F_y \Delta y + F_z \Delta z. \quad (14.28)$$

Il contributo di $F_x \Delta x + F_y \Delta y$ deve essere mediamente nullo, perché non c'è variazione di energia se $z = \text{cost}$, quindi possiamo scrivere che

³In effetti le stesse considerazioni si applicano a una qualsiasi delle tre componenti del momento angolare.

$$\mathbf{F} = (0, 0, F_z) = \left(0, 0, \mu_z \frac{\Delta B_z}{\Delta z}\right). \quad (14.29)$$

Se quindi $\mathbf{B}$ non è costante lungo la direzione z , ma varia, sull'atomo si produce una forza lungo la stessa direzione. Dal momento che la direzione z è arbitraria lo stesso vale per le altre componenti, quindi, in generale

$$\mathbf{F} = (F_x, F_y, F_z) = \left(\mu_x \frac{\Delta B_x}{\Delta x}, \mu_y \frac{\Delta B_y}{\Delta y}, \mu_z \frac{\Delta B_z}{\Delta z}\right). \quad (14.30)$$

Si potrebbe provare quindi a realizzare un campo magnetico per cui B_x e B_y risultano costanti, mentre B_z risulta variabile con z , pertanto un atomo (classico) con momento magnetico qualunque subirà forze dirette lungo z che potranno avere intensità variabile da 0 a $\mu \Delta B_z / \Delta z$, visto che la componente μ_z può assumere valori compresi tra 0 e μ e un verso positivo o negativo. Questo, purtroppo non è possibile a causa delle equazioni di Maxwell. Secondo una di queste equazioni la divergenza di $\mathbf{B}$ è nulla (il che equivale a dire che il flusso attraverso una qualunque superficie chiusa di questo campo è nullo). Questo implica che

$$\frac{\Delta B_x}{\Delta x} + \frac{\Delta B_y}{\Delta y} + \frac{\Delta B_z}{\Delta z} = 0 \quad (14.31)$$

quindi almeno un'altra componente di B deve variare in modo che questa somma sia nulla. Se ammettiamo che B non vari lungo x dev'essere

$$\frac{\Delta B_y}{\Delta y} = -\frac{\Delta B_z}{\Delta z}. \quad (14.32)$$

La traiettoria di questi atomi con velocità iniziale lungo l'asse x in questa regione è curva e dopo aver percorso una data lunghezza nella direzione x raggiungono una coordinata z che è un numero qualunque compreso tra $-z_0$ e $+z_0$. Se si esegue l'esperimento con atomi di varia specie, dal momento che questi non si comportano come atomi *classici*, ma quantistici, si vede che la coordinata raggiunta

dopo aver percorso una distanza x non è un valore qualunque compreso tra due estremi, ma un insieme di valori discreti: il momento magnetico dunque è quantizzato come ci si aspetta per il momento angolare cui dovrebbe essere proporzionale.

Moto di momenti magnetici classici in campi non uniformi

Se un momento magnetico $\boldsymbol{\mu} = (0, 0, \mu_z)$ orientato lungo l'asse z si muove in campo magnetico non uniforme subisce una forza $\mathbf{F} \simeq (0, 0, F_z) = \mu_z (0, 0, \frac{\Delta B_z}{\Delta z})$. Per la Legge di Newton $\mathbf{F} = m\mathbf{a}$ e il moto è dunque accelerato con accelerazione lungo l'asse z . Se $\Delta B_z / \Delta z = \text{cost}$, il moto è uniformemente accelerato. Supponiamo che la particella con momento magnetico $\boldsymbol{\mu}$ si muova inizialmente con velocità $\mathbf{v}_0 = (v_0, 0, 0)$ partendo dal punto di coordinate $x_0 = (0, 0, 0)$. L'equazione del moto della particella è dunque

$$\begin{cases} x = v_0 t \\ y = 0 \\ z = \frac{1}{2} \mu_z \frac{\Delta B_z}{\Delta z} t^2 \end{cases} \quad (14.33)$$

e dopo aver percorso una distanza L lungo l'asse x si ha $t = L/v_0$ e dunque

$$z = \frac{1}{2} \mu_z \frac{\Delta B_z}{\Delta z} \frac{L^2}{v_0^2} \quad (14.34)$$

Naturalmente, se il momento magnetico dell'atomo fosse nullo la traiettoria sarebbe comunque rettilinea perché l'accelerazione sarebbe nulla e tutti gli atomi andrebbero a urtare nello stesso punto un eventuale bersaglio posto lungo l'asse x .

Questo esperimento fu eseguito da **Otto Stern** e **Walther Gerlach** nel 1922 [14]. Non senza sorpresa, eseguendo l'esperimento con atomi di argento, per i quali $L = 0$, si vide che questi atomi vengono in effetti deviati: alcuni verso l'alto, altri verso il basso, e colpiscono il bersaglio in due punti distinti. Non esistono atomi che attraversano la regione di campo magnetico e sono deviati di una quantità interme-

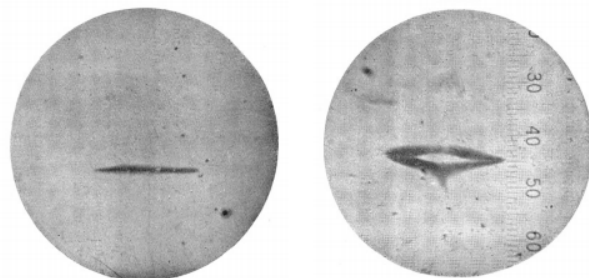


Figura 14.5 Immagini realizzate da Stern e Gerlach inviando fasci di atomi d'argento su una lastra fotografica. A sinistra si vede l'immagine prodotta in assenza di campo magnetico. L'asse z è verticale e quello y orizzontale; l'asse x entra nel piano dell'immagine. In questo caso gli atomi si muovono di moto rettilineo uniforme lungo x e hanno coordinata z pari a zero e coordinata y variabile (il fascio era allargato in questa direzione). Accendendo il campo magnetico (a destra) la riga si separa in due archi ben distinti (l'intensità del campo e della sua variazione lungo z diminuisce progressivamente allontanandosi dal centro).

dia tra queste due. In particolare nessun atomo percorreva la traiettoria attesa: quella rettilinea. Se ne deve concludere che gli elettroni di questi atomi (e quindi tutti gli elettroni) devono possedere un momento magnetico **intrinseco** e dunque un momento angolare **intrinseco**, che non dipende dal fatto di avere una velocità relativa rispetto a qualcos'altro. Questo momento angolare intrinseco è chiamato **spin**.

In alcuni testi la presenza dello spin è associata a una sorta di rotazione dell'elettrone su sé stesso, in analogia a quel che succede per un pianeta che ruota attorno al Sole e attorno al proprio asse: quel pianeta ha un momento angolare $L = mvr$ rispetto al Sole, da cui dista r , e un momento angolare intrinse-

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.6 Animazione dell'esperimento
di Stern–Gerlach. Estratto
da www.vulgarisation.fr

co che dipende dal fatto di ruotare su sé stesso. Ma questa immagine è completamente sbagliata! Non si può parlare di rotazione dell'elettrone su sé stesso, per molte ragioni. Quanto meno, un'eventuale rotazione sarebbe inosservabile per i principi della meccanica quantistica e nella fisica non possiamo utilizzare un concetto privo di senso. Cercare di farci un'immagine classica di un fenomeno puramente quantistico di solito è il modo peggiore di procedere. Il modo giusto di afferrare il concetto è quello di pensare allo spin come a una grandezza fisica che si può misurare, indipendentemente da come possiamo immaginarci l'elettrone, che ha le dimensioni e il comportamento di una grandezza classica come il momento angolare dovuto alla rotazione, ma che non ha niente a che vedere con questa!

Lo spin è un momento angolare quantistico e quindi i suoi valori differiscono di ± 1 l'uno dall'altro. Avendo osservato solo due valori e l'assenza del valore nullo ne concludiamo che lo spin di un elettrone può assumere solo due valori diversi da zero che non possono che essere, in unità di $\hbar$, $+\frac{1}{2}$ e $-\frac{1}{2}$.

Un elettrone in un atomo dunque possiede un momento angolare totale J che è la somma del momento angolare orbitale L e del suo spin S : $J = L + S$. Anche J è un momento angolare e deve essere anche lui quantizzato, dunque la somma non si può eseguire banalmente come una comune somma. Il massimo valore del momento angolare totale lo si ha quando L e S sono tra loro paralleli e hanno lo stesso verso. In questo caso le loro terze componenti si sommano e il momento angolare totale effettivamente vale $J = L + S$. Ma L e S potrebbero essere discordi e in questo caso le loro terze componenti si sottrarrebbero dando luogo a un momento angolare totale di modulo pari a $J = |L - S|$. Essendo quantizzato, il momento angolare totale di un elettrone in un ato-

mo può quindi assumere tutti i valori compresi tra $|L - S|$ e $L + S$, a passi di 1.

Se quindi abbiamo $L = 0$ esiste un solo valore possibile per $J = 1/2$, con due possibili orientazioni: $-1/2$ e $+1/2$. Per $L = 1$ i possibili valori di J vanno da $J = |L - S| = 1/2$ a $J = L + S = 3/2$. I possibili valori della terza componente sono, in un caso $-1/2$ e $+1/2$, nell'altro $-3/2$, $-1/2$, $+1/2$, $+3/2$. Per $L = 2$ i possibili valori di J sono compresi tra $J = |L - S| = 3/2$ e $J = L + S = 5/2$, e sono quindi tre: $3/2$, $4/2 = 2$ e $5/2$. I rispettivi valori delle terze componenti sono, nel caso di $J = 3/2$, $-3/2$, $-1/2$, $+1/2$, $+3/2$; nel caso $J = 2$, -2 , -1 , 0 , 1 e 2 ; nel caso $J = 5/2$, $-5/2$, $-3/2$, $-1/2$, $1/2$, $3/2$, $5/2$.

14.9 Il Principio di Pauli

Sebbene nessuno lo faccia mai notare, in meccanica classica esiste un **Principio di esclusione** molto evidente: due punti materiali non possono avere contemporaneamente lo stesso stato. Lo stato di un punto materiale infatti è completamente determinato quando se ne conoscano contemporaneamente posizione e velocità. È del tutto evidente che due oggetti non possono stare esattamente nello stesso punto dello spazio mantenendo la stessa velocità: se vengono a trovarsi nello stesso punto dello spazio si urtano e devono quindi avere velocità diverse e se hanno velocità diverse, all'istante immediatamente successivo avranno di certo coordinate diverse.

La teoria della relatività impone di cambiare la definizione di stato perché la velocità non è una buona misura dello stato di una particella. Lo stato dev'essere determinato dalla posizione e dalla quantità di moto della particella. La natura del principio però resta.

Con l'avvento della meccanica quantistica anche la definizione relativistica di stato vacilla: posizione e quantità di moto non si possono mai conoscere contemporaneamente e dunque non sono una buona definizione di stato. Lo stato di un oggetto è costituito dall'insieme delle variabili che permettono di determinarne l'evoluzione futura. Nel caso quantistico, due grandezze che si possono misurare

contemporaneamente e che permettono di predire lo stato a un tempo successivo sono, per esempio, l'energia e il momento angolare totale (insieme alla sua terza componente). Se si accetta questo punto di vista (e non si può non accettare dal momento che tutte le evidenze sperimentali vanno in questa direzione), quello che è spesso considerato un Principio astruso e incomprensibile come il **Principio di esclusione di Pauli**⁴ diventa del tutto naturale, trattandosi né più né meno che dello stesso principio di esclusione vigente nella meccanica classica e del tutto comprensibile.

14.9.1 La chimica

Se l'energia di un elettrone in un atomo è limitata a un certo valore (che dipende dal livello occupato n), lo è anche la sua distanza media da questo (la distanza non si può definire, ma il suo valor medio sí). Di conseguenza è limitato anche il momento angolare $L = mvr$. Questo significa che il massimo momento angolare che un elettrone può avere in un atomo, in unità di $\hbar$ è limitato a $\ell \leq \ell_0$, dove ℓ_0 deve dipendere da n perché il limite esiste in quanto conseguenza del fatto che l'energia è limitata. Si può dimostrare che $\ell \leq n - 1$. Gli elettroni con energia minima (quelli con $n = 1$) possono quindi avere soltanto $L = 0$, mentre quelli di energia superiore (con $n = 2$) possono avere sia $L = 0$ che $L = 1$, e così via.

La **tavola periodica** di **Mendeleev** (Fig. 14.7) si spiega molto bene alla luce di questa teoria. La posizione nella tavola di ogni elemento ne fornisce la **valenza**, che è sempre un numero intero. Alcuni elementi si combinano in modo tale che la somma delle rispettive valenze faccia 2, altri in modo che faccia 8, altri ancora 10 e così via. Da dove vengono questi numeri magici?

È molto semplice: un atomo d'idrogeno ha un solo elettrone che, trovandosi nello stato di energia più bassa possibile potrà avere solo $L = 0$. Per il Principio di esclusione di Pauli solo un altro elettrone può trovare posto in questo stesso stato di energia

e momento angolare: quello con lo spin rivolto nella direzione opposta. Per questa ragione l'idrogeno è **monovalente**. Non c'è più posto per altri elettroni in questo stesso stato. In effetti se un elettrone avesse un momento angolare maggiore, avrebbe anche un'energia più alta.

L'elio ha due elettroni che possono stare nello stato con $L = 0$, che a questo punto è saturo. L'elio è infatti un gas nobile che non reagisce chimicamente con altri elementi. Aggiungendo un elettrone si ha il litio. Poiché nello stato di minima energia non c'è posto per tre elettroni, ma per due, il terzo deve trovarsi in uno stato di energia più elevata con L però che può ancora essere pari a 0. In questo stato c'è di nuovo posto per un altro elettrone e il litio è monovalente come l'idrogeno. Il berillio ha quattro elettroni che saturano lo stato di energia più alto con $L = 0$, ma questo stato può avere anche elettroni con $L = 1$. Quindi il berillio non è un gas nobile, ma ha valenza 2 e si lega a elementi con valenza 6 (o due con valenza 3). Sono infatti 8 i possibili stati di momento angolare per elettroni di energia pari a quella dell'elettrone più energetico del berillio: due per $L = 0$ e sei per $L = 1$. In questo stato c'è posto fino a 8 elettroni e fino a quando lo stato non è saturo è sempre possibile che un altro elettrone possa piazzarsi. Dopo il berillio vengono infatti il boro, il carbonio, l'azoto, l'ossigeno, il fluoro e il neon, per un totale di otto possibili elementi in questa riga. Tutti questi elementi hanno in comune l'energia massima dei propri elettroni e un numero di elettroni in questo stato energetico variabile tra 1 e 8. Il neon, che ne ha otto, è un gas nobile. Il fluoro ne ha sette, dunque può accettare un elettrone da parte di un altro atomo e per questo è monovalente. La riga successiva comincia col sodio, che è monovalente perché ha un solo elettrone nello stato più esterno. Il penultimo elemento di questa riga è il cloro, che ne ha sette. Avvicinando un atomo di cloro a uno di sodio, l'elettrone del sodio non può più dire di appartenere al sodio: in effetti si troverà a essere **distribuito** attorno al nucleo di cloro e di sodio al tempo stesso. Lo stesso vale per gli elettroni del cloro e così si forma il legame chimico. La riga successiva contiene molti più elementi: per la

⁴Dal nome del fisico **Wolfgang Pauli**.

precisione ce ne sono dieci in piú. In effetti gli elettroni piú energetici degli atomi della riga del sodio, avendo $n = 3$, dovrebbero avere $L = 0, 1$ o 2 . Ci aspetteremmo quindi non otto, ma $8 + 15 = 23$ elementi diversi. Il fatto è che la configurazione nella quale l'energia è piú alta, ma il momento angolare è piú basso risulta favorita rispetto a quella con energia piú bassa e momento angolare alto, perciò gli elementi della terza riga hanno gli elettroni negli stati $L = 0, 1$, ma non $L = 2$. Questa configurazione è sfavorita rispetto a quella in cui un elemento ha ancora $L = 0$, ma energia piú alta, corrispondente al caso del potassio. Gli elementi con $Z = 21$ fino a $Z = 30$, che sono 10, sono proprio quelli che hanno energia piú bassa, ma momento angolare piú alto, pari a $L = 2$. Combinandosi con lo spin, questo fa $J = 3/2$, $J = 2$ o $J = 5/2$. Ancora una volta gli stati con $J = 3/2$ e $J = 5/2$ sono favoriti rispetto a quelli con $J = 2$. Questi stati sono proprio 10 e per questo dopo lo zinco ($Z = 30$) viene il gallio, a $Z = 31$ con $L = 1$ ed energia maggiore.

Al crescere dell'energia e di L le configurazioni possibili si complicano, ma resta una certa regolarità per la quale comunque sono favorite sempre le configurazioni con $L = 0$ e $L = 1$, che danno sempre otto possibili combinazioni. Talvolta le combinazioni possono essere di piú, ma non sono così frequenti: per questo la maggior parte degli elementi reagisce con altri in modo tale che la somma delle valenze faccia per lo piú otto: la somma delle valenze non è altro che il numero massimo di elettroni che trovano posto nello stato di massima energia possibile per quegli elementi.

The Periodic Table of the Elements

Legend:

- alkali metals
- alkaline earth metals
- transition metals
- lanthanoids
- actinoids
- metalloids
- nonmetals
- halogens
- noble gases
- unknown elements
- radioactive elements have masses in parenthesis

Iron (Fe) Detailed View:

- atomic mass: 55.845
- atomic number: 26
- chemical symbol: Fe
- name: Iron
- electron configuration: [Ar] 3d⁶ 4s²
- oxidation states: +2, +3

Notes:

- as of yet, elements 113-118 have no official name designated by the IUPAC.
- 1 kJ/mol = 96.485 eV.
- all elements are implied to have an oxidation state of zero.

Figura 14.7 La tavola periodica degli elementi chimici.

14.9.2 Semiconduttori

Se pensate che la meccanica quantistica non abbia alcun interesse pratico vi sbagliate di grosso! Infatti questo tipo di fisica è alla base del funzionamento dei semiconduttori con i quali sono realizzati tutti gli apparati elettronici: dal telecomando dell'auto al tablet.

Un atomo di silicio è un atomo che nello stato energetico più elevato ospita quattro elettroni. C'è posto dunque per altri quattro. I quattro elettroni presenti hanno tutti la stessa energia, ma momento angolare diverso. La forza con la quale il nucleo *trattiene* gli elettroni attorno a sé è abbastanza forte e un cristallo di silicio si comporta in genere come un **isolante**: applicando una differenza di potenziale gli elettroni non si muovono, perché trattenuti dalle forze Coulombiane che esercitano i nuclei del materiale, e la **corrente** non passa. Se si scaldasse il silicio a una temperatura sufficientemente alta oppure fosse illuminato da una radiazione ultravioletta, gli elettroni acquisterebbero energia sufficiente a portarsi nello stato di energia più elevato. Sarebbero così meno legati al nucleo e più liberi di muoversi. Applicando una differenza di potenziale quindi si osserverebbe il passaggio di corrente. Da isolante, il silicio diventa conduttore.

In un cristallo di silicio gli elettroni però non sono distribuiti nello stesso modo che in un atomo. Per capire come mai, supponiamo che il silicio abbia un solo elettrone e che se ne possa trovare uno solo in un determinato stato di energia: è un'ipotesi semplificativa che tuttavia non cambia la sostanza del meccanismo. Un atomo sarebbe dunque formato da un nucleo circondato dal suo elettrone. Due atomi diversi di silicio avrebbero i due elettroni nello stesso stato, di energia

$$U = -kZ \frac{e^2}{r} \quad (14.35)$$

dove r è la distanza media tra nucleo ed elettrone. Quando i due atomi si avvicinano, però, l'elettrone dell'uno sente, oltre al campo Coulombiano del proprio nucleo, anche quello dell'altro e quello pro-

dotto dall'altro elettrone (che ha segno opposto). L'energia dell'elettrone diventa quindi

$$U = -kZ \frac{e^2}{r_1} - kZ \frac{e^2}{r_2} + k \frac{e^2}{r_3} \quad (14.36)$$

dove r_i sono le distanze, rispettivamente, dal nucleo, dall'altro nucleo e dall'altro elettrone. È evidente che l'energia di ogni elettrone cambia rispetto alla configurazione precedente e dal momento che anche in questo caso l'energia dell'elettrone sarà quantizzata avremo degli stati di energia discreti

$$E_n = -\frac{A}{n^2}, \quad (14.37)$$

con A costante. Non essendo possibile per due elettroni stare nello stesso stato, ognuno dovrà piazzarsi in uno stato diverso: un elettrone assume l'energia più bassa possibile (per $n = 1$), mentre l'altro quella immediatamente superiore (per $n = 2$). Nessuno dei due elettroni potrà più dire a quale nucleo appartiene: di fatto ogni elettrone *circonda* tutti e due i nuclei.

Se avviciniamo un terzo atomo accadrà qualcosa di simile: i tre elettroni si troveranno ciascuno in un livello energetico diverso, separato dagli altri da un intervallo di energie proibite. È facile capire che all'aumentare degli elettroni la perturbazione che si provoca sull'energia di ciascuno diventa via via più piccola e quindi la *distanza* tra i livelli energetici diminuisce sempre di più. In altre parole, con due nuclei i primi due livelli sono separati di

$$\Delta E_2 = -A_2 \left(\frac{1}{4} - 1 \right), \quad (14.38)$$

mentre con N nuclei la separazione tra i livelli diventa

$$\Delta E_N = -A_N \left(\frac{1}{4} - 1 \right), \quad (14.39)$$

dove $A_N \ll A_2$ e A_N tende a zero man mano che aumenta N . In un cristallo ci sono dell'ordine di $N_A \simeq 6 \times 10^{23}$ (un **numero di Avogadro**) atomi, perciò i livelli energetici occupati dagli elettroni, pur discreti, sono talmente fitti da costituire una vera e propria **banda** continua. In certe condizioni

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.8 Cristalli di tipo p e n conducono l'elettricità in modi diversi: nei primi sono le lacune (positive) a condurre l'elettricità; nei secondi sono gli elettroni (negativi) [<https://www.youtube.com/watch?v=reaizv3jjCY>].

è possibile che esista una banda di energie permesse separata da un'altra banda di energia permesse da un intervallo di energie proibite. Se gli elettroni si trovano tutti nella prima banda (detta **banda di valenza**) il materiale si comporta come un isolante, mentre se alcuni elettroni si trovano nella banda superiore (detta **banda di conduzione**) il materiale diventa conduttore. Il silicio naturale avrebbe tutti gli elettroni nella banda di valenza e dunque sarebbe un isolante.

Se però inserisco nel reticolo cristallino, di tanto in tanto, alcuni atomi di Arsenico (che di elettroni ne ha cinque), quattro di questi trovano posto nella banda di valenza e uno in quella di conduzione (o nelle immediate vicinanze). Il cristallo così formato (che si dice **drogato di tipo n**) è quindi un conduttore a temperatura ambiente.

Se, al contrario, s'inseriscono nel reticolo atomi di Boro (trivalente), i tre elettroni che possiede nello stato di energia più alta finiscono tutti nella banda di valenza (in realtà appena un po' più su, ma la sostanza non cambia) dove ci sarebbe posto per un elettrone che non c'è. Esiste dunque un livello energetico tra i tanti libero da elettroni. Il cristallo si dice **drogato di tipo p**. Il motivo è presto spiegato: se si applica una differenza di potenziale ai capi di questo cristallo gli elettroni nella banda di valenza acquistano un po' di energia, che non è sufficiente a farli passare nella banda di conduzione, ma è abbastanza per far passare un elettrone dal livello, diciamo E_n al livello E_{n+1} , infinitamente vicino. Questo può avvenire solo se il livello di ener-

gia E_{n+1} è vuoto, e se avviene quello di energia E_n si svuota. Il processo si può reiterare a piacere e la conseguenza netta è che la **mancanza** di un elettrone (quella che si chiama una **lacuna**) passa da un livello a un altro in modo esattamente opposto a quel che farebbe un elettrone. È come se nel cristallo fosse presente una carica positiva che, sotto l'azione della differenza di potenziale applicata, si sposta nel cristallo in modo opposto a come fa un elettrone. Per questo il cristallo si dice di tipo p : le cariche che conducono la corrente sono **positive**⁵.

14.9.3 Il diodo

Se si fa crescere un cristallo drogato di tipo p sopra un cristallo drogato di tipo n , all'interfaccia tra i due alcune cariche negative sono libere di muoversi da un lato, mentre dall'altro esistono alcuni stati di bassa energia liberi da elettroni. Per un elettrone libero di muoversi nel cristallo n è quindi molto facile occupare il livello vuoto presente nel cristallo di tipo p , ma una volta *caduto* in quel livello diventa sostanzialmente immobile non essendo più possibile, per lui, cambiare energia dal momento che tutti gli stati di energia vicini sono occupati. Viceversa, nel cristallo di tipo n si produce una lacuna nella banda di conduzione, perché quello stato è stato abbandonato da un elettrone.

Di fatto si produce una condizione per cui all'interfaccia tra un cristallo e l'altro, si trova uno strato di cariche negative in eccesso nel cristallo p e un eccesso di cariche positive (lacune) nel cristallo n (il cristallo di tipo n è elettricamente neutro perché dove si trovano gli atomi di arsenico ci sono 5 protoni e altrettanti elettroni: se un elettrone abbandona il cristallo per passare nell'altro si crea uno squilibrio di cariche).

Il dispositivo così formato è un **diodo**. I diodi funzionano come valvole idrauliche: fanno passare la corrente in un verso, ma non nell'altro. Se infatti si collega il polo positivo di una pila al cristallo di tipo n e quello negativo al cristallo di tipo p gli elet-

⁵Notate che in questo caso la corrente è sempre dovuta al moto degli elettroni, che però sono nella banda di valenza.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.9 Un diodo funziona come una valvola idraulica: lascia passare la corrente che scorre in un verso, ma non quella che scorre nell'altro [<https://www.youtube.com/watch?v=y2htodNi8xI>].

troni nel cristallo tenderanno a muoversi nel verso che va dal cristallo p al cristallo n . Nel far questo troveranno all'interfaccia uno strato di elettroni che li respingerà impedendo loro di passare, da un lato; dall'altro lo strato di cariche positive tenderà a trattenere gli elettroni. Se invece il diodo si **polarizza** al contrario gli elettroni possono muoversi dal cristallo n al cristallo p : questo moto è addirittura favorito dallo strato di cariche positive presenti all'interfaccia e quindi il cristallo conduce.

Con uno strumento così si possono fare molte cose: ad esempio i sensori delle fotocamere digitali. Se si espone un diodo alla luce, per [effetto fotoelettrico](#) alcuni elettroni si liberano e passano nella banda di conduzione. Ma questi ricadono presto nella banda di valenza e comunque si muovono in tutte le direzioni possibili rendendo nulla la corrente fotoelettrica media. Se però il diodo è polarizzato in modo da non condurre corrente, gli elettroni del cristallo non si spostano, ma quelli che sono stati portati nella banda di conduzione dall'energia della radiazione luminosa, si muovono tutti coerentemente dal polo negativo a quello positivo della pila. In questo modo si produce una sia pur debole corrente proporzionale all'intensità della luce che ha colpito il **pixel** di silicio. Le fotocamere digitali dunque possono funzionare solo grazie alla meccanica quantistica che permette l'effetto fotoelettrico e la formazione delle bande: se l'energia degli elettroni in un atomo potesse assumere ogni valore possibile i diodi non si potrebbero costruire.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 14.10 Il funzionamento di un transistor è analogo a quello di uno *sciacquone* [<https://www.youtube.com/watch?v=IQJnGhsWfWc>].

14.9.4 Il transistor

Con tre cristalli di tipo alternato (*npn* o *pnp*) si fanno, invece, i **transistor**, che funzionano come **amplificatori** di corrente. I tre semiconduttori sono collegati ad altrettanti conduttori detti, rispettivamente **emettitore**, **base** e **collettore**. I transistor amplificano la corrente che s'inietta nella loro base. Dal momento che la carica elettrica si conserva non è evidentemente possibile moltiplicare davvero il numero di cariche che circolano nel transistor, ma si può creare l'*illusione* che questo sia possibile.

Possiamo immaginare il funzionamento di un transistor come quello di una coppia di valvole idrauliche montate al contrario⁶ (vedi il Filmato 14.10) e il suo funzionamento come quello di uno **sciacquone**. L'emettitore di un transistor si collega a una pila che non è altro se non una riserva di cariche, analogamente alla cassetta di scarico di un bagno che serve per immagazzinare acqua. Collegando il dispositivo di scarico a un secchiello (che funge da base del transistor) possiamo provocare l'attivazione introducendo acqua nel secchiello (dovete pensare a uno di quei sciacquoni con la catena che ormai non s'usano più). Il peso dell'acqua provoca l'apertura della valvola che fa cadere l'acqua dalla cassetta/emettitore al water/collettore.

A fronte di un modesto quantitativo d'acqua nel secchiello/base, vediamo scorrere tanta acqua nel water/collettore. È come se avessimo moltiplicato la quantità d'acqua versata per un fattore, ma naturalmente l'acqua che scorre nel water è quella pre-

⁶In un certo senso si potrebbe immaginare come una coppia di diodi montati al contrario l'uno rispetto all'altro, ma in questo caso la presenza del conduttore tra l'uno e l'altro impedirebbe al sistema il funzionamento come transistor.

levata dalla cassetta che non ha niente a che vedere con quella introdotta nel dispositivo. Come la cassetta, anche la pila di un transistor infatti si scarica con l'uso.

I transistor sono l'elemento principale con il quale si costruiscono i circuiti integrati che si trovano, a milioni, in ogni dispositivo elettronico. Anche il transistor è un dispositivo quantistico perché anche il suo funzionamento dipende dalla struttura a bande caratteristica di questo tipo di fisica. La meccanica quantistica, insomma, è molto meno esotica di quanto si creda!

14.10 L'equazione di Schrödinger

Tutte le evidenze sperimentali suggeriscono che gli oggetti di cui è costituito l'Universo sono descrivibili matematicamente in termini di onde e di particelle al tempo stesso. Inoltre la fisica può solo aver a che fare con grandezze **misurabili**. Il processo di misura dunque è parte essenziale dell'evoluzione di un sistema fisico che deve essere descritto in termini dei risultati di una misura.

Una delle grandezze fisiche che si possono misurare su uno stato ψ che descrive un sistema fisico è l'energia. La misura consiste nell'interazione con uno strumento che dà come risultato il valore dell'energia dello stato E . Possiamo rappresentare l'interazione con un **operatore** matematico H **moltiplicato** per lo stato stesso: $H\psi$. Il risultato dell'applicazione di quest'operatore non può che essere il risultato della misura E che, per ragioni dimensionali, deve includere lo stato ψ . Possiamo dunque scrivere che

$$H\psi = E\psi \quad (14.40)$$

dove E è un numero reale, mentre H è qualcosa che si applica a ψ e dà come risultato E moltiplicata per ψ . Se rappresentiamo lo stato di una particella come un'onda possiamo scrivere che

$$\psi = \psi(x) = A \sin(kx + \phi) \quad (14.41)$$

dove k è il numero d'onda e ϕ una fase arbitraria. Il numero d'onda è legato alla lunghezza d'onda dalla relazione

$$k = \frac{2\pi}{\lambda} \quad (14.42)$$

che in meccanica quantistica e in unità naturali altro non è che la quantità di moto della particella p (vedi equazione (14.18)). Possiamo quindi scrivere che, in assenza d'interazioni e per velocità non relativistiche,

$$H A \sin(kx + \phi) = E A \sin(kx + \phi) = \frac{p^2}{2m} A \sin(kx + \phi). \quad (14.43)$$

Applicare l'operatore H a ψ dunque deve risultare nella moltiplicazione di questo per $p^2/2m = k^2/2m$, cioè $H = p^2/2m$. Le dimensioni fisiche sono quelle che ci si aspettano (quelle di un'energia) ed è abbastanza naturale estendere la definizione di H nel caso in cui sia presente un campo di forze il cui potenziale sia V come

$$H = \frac{p^2}{2m} + V, \quad (14.44)$$

per la quale avremo che

$$\left(\frac{p^2}{2m} + V\right)\psi = E\psi \quad (14.45)$$

nota come **equazione di Schrödinger**⁷. Osserviamo che la derivata di ψ rispetto a x vale

$$\frac{d\psi}{dx} = kA \cos(kx + \phi) \quad (14.46)$$

e quindi la derivata seconda

$$\frac{d^2\psi}{dx^2} = -k^2 A \sin(kx + \phi), \quad (14.47)$$

pertanto possiamo dire che la quantità di moto in fisica quantistica altro non è se non i volte la derivata rispetto a x dello stato dal momento che dividendo tutto per $-2m$ abbiamo

⁷Dal nome del fisico **Erwin Schrödinger** che la propone.

$$-\frac{1}{2m} \frac{d^2\psi}{dx^2} = \frac{k^2}{2m} \psi \quad (14.48)$$

che riproduce l'equazione (14.43) se ammettiamo che

$$p = i \frac{d}{dx}. \quad (14.49)$$

Il suo quadrato infatti è

$$p^2 = \left(i \frac{d}{dx}\right)^2 = -\frac{d^2}{dx^2} \quad (14.50)$$

e si può pensare ad H non come a un semplice numero che moltiplica lo stato, ma come un'operazione che consiste nel calcolare la derivata seconda dello stato rispetto alle coordinate, moltiplicata per $1/2m$.

Nell'equazione di Schrödinger lo stato ψ , che è una funzione di x è rappresentato come una **funzione d'onda**. Se si integra su tutto l'asse x il secondo membro dell'equazione di Schrödinger si ha

$$\int_{-\infty}^{\infty} E\psi(x)dx. \quad (14.51)$$

Integrare su tutto l'asse x significa misurare tutte le possibili energie che la particella rappresentata da $\psi(x)$ può avere nei diversi punti dello spazio e sommarle tra loro, pesandole con ψ . Se si impone che

$$\int_{-\infty}^{\infty} \psi(x)dx = 1 \quad (14.52)$$

l'integrale sopra non è altro che il valor medio dell'energia quando la funzione di distribuzione di probabilità dei suoi valori è $\psi(x)$. Di conseguenza il modulo quadro di $\psi(x)$, $|\psi(x)|^2$ si può interpretare come la probabilità di trovare la particella nel punto x . Quando dunque scriviamo lo stato di una particella come una funzione d'onda diciamo che la particella, fino a che non si esegue una misura, si trova ovunque nello spazio con probabilità diverse (che possono anche essere nulle). Questa interpretazione della funzione d'onda è coerente con quanto visto sopra. Consideriamo, per esempio, l'esperimento della doppia fenditura nel quale s'invia un

fascio di elettroni su uno schermo con due sottili fenditure. Prima di arrivare allo schermo gli elettroni non sono localizzati (se la sorgente è molto lontana) e quindi la loro funzione d'onda si estende in tutto lo spazio: ogni singolo elettrone è **ovunque** dietro lo schermo. Giungendo sullo schermo, quest'ultimo di fatto *esegue una misura* della posizione degli elettroni. Al di là dello schermo la funzione d'onda di un elettrone è di fatto nulla dappertutto tranne che in corrispondenza delle fenditure dove la probabilità di trovare l'elettrone è diversa da zero. Ma attenzione! Non si tratta semplicemente del fatto che **noi** non conosciamo la sua posizione: il fatto è che la sua posizione non è determinata e si trova con probabilità $1/2$ e contemporaneamente in corrispondenza dell'una e dell'altra fenditura. Solo così è possibile che, propagandosi al di là di queste, le due funzioni d'onda, che a questo punto hanno fronti semicircolari, interferiscono tra loro e, giungendo su un altro schermo, producono una figura d'interferenza. Rivelare gli elettroni sullo schermo equivale a misurare la loro posizione: ogni elettrone risulterà essere a una certa posizione con una probabilità che dipende da come hanno interferito le onde prodotte dalle due fenditure, quindi di fatto la figura d'interferenza è una rappresentazione del modulo quadro di ψ .

È utile rappresentare la funzione d'onda sfruttando la proprietà dell'equazione di Schrödinger di essere lineare. Se infatti $A \sin(kx + \phi)$ è una soluzione lo è anche $B \cos(kx + \phi)$ e dunque lo è anche una combinazione lineare delle due:

$$\psi = A \sin(kx + \phi) + B \cos(kx + \phi). \quad (14.53)$$

Il modulo quadro di ψ deve valere 1 quindi A e B si devono scegliere in modo tale da soddisfare questa condizione. Un modo semplice d'imporre questa condizione consiste nello scrivere la soluzione generale come

$$\psi = \sin(kx + \phi) + i \cos(kx + \phi). \quad (14.54)$$

Il modulo quadro di ψ è $\psi\psi^*$, dove ψ^* rappresenta il complesso coniugato di ψ :

$$\psi^* = \sin(kx + \phi) - i \cos(kx + \phi). \quad (14.55)$$

In questo modo la condizione di normalizzazione (14.52) è automaticamente soddisfatta. Usando infine le [formule di Eulero](#) possiamo scrivere la funzione d'onda come

$$\psi = e^{i(kx + \phi)} \quad (14.56)$$

Allo stesso modo si spiega un altro fenomeno curioso: un elettrone libero ha spin che può esistere in due stati: $+\frac{1}{2}$ e $-\frac{1}{2}$. Se misuriamo lo spin di un elettrone usando un apparato tipo quello di Stern e Gerlach orientato per esempio in modo l'asse z sia verticale, possiamo dire se ha spin $+\frac{1}{2}$ o spin $-\frac{1}{2}$, mentre prima della misura dobbiamo ammettere che si trovi in uno stato che si può rappresentare come

$$\psi = \frac{1}{\sqrt{2}} \left| +\frac{1}{2} \right\rangle + \frac{1}{\sqrt{2}} \left| -\frac{1}{2} \right\rangle \quad (14.57)$$

dove i fattori $1/\sqrt{2}$ servono a garantire che la probabilità di trovare l'elettrone in uno stato o nell'altro sia 1. Supponendo d'aver misurato lo stato $+\frac{1}{2}$ lo stato diventa

$$\psi = \left| +\frac{1}{2} \right\rangle. \quad (14.58)$$

Se ora l'elettrone continua a propagarsi ed eseguiamo nuovamente una misura non possiamo che ritrovare lo stesso stato. Ma se l'apparato di Stern e Gerlach è ruotato di 90° rispetto al primo (lo disponiamo quindi in modo da avere l'asse z orizzontale) misuriamo una grandezza fisica diversa: la componente dello spin in una direzione diversa. Dal momento che questa prima non era mai stata misurata l'elettrone si trova in uno stato composto dalla sovrapposizione dei due stati e quindi nuovamente la sua funzione d'onda diventa

$$\psi = \frac{1}{\sqrt{2}} \left| +\frac{1}{2} \right\rangle + \frac{1}{\sqrt{2}} \left| -\frac{1}{2} \right\rangle \quad (14.59)$$

dove ora gli spin sono misurati in una direzione diversa. Classicamente dovremmo aspettarci un valore

nullo per lo spin in direzione perpendicolare a quella misurata in precedenza, ma questo non è possibile perché il valore zero non è ammesso tra i possibili risultati di una misura. Misureremo quindi, con uguale probabilità, uno spin orientato a destra o a sinistra, benché immediatamente prima abbiamo accertato che quest'ultimo era orientato in su o in giù.

Un'onda che si propaga, oltre a essere funzione dello spazio, è anche funzione del tempo perciò possiamo scrivere $\psi = \psi(t)$ e rappresentare ψ come

$$\psi = \psi(t) = e^{-i(\omega t + \phi)} \quad (14.60)$$

Facendo la derivata della funzione d'onda rispetto al tempo si trova

$$\frac{d\psi}{dt} = -i\omega e^{-i(\omega t + \phi)} = -2i\pi\nu e^{-i(\omega t + \phi)}. \quad (14.61)$$

Dividendo per 2π e moltiplicando per $i\hbar$, ricordando che $\hbar = h/2\pi$ e che $E = h\nu$, si ottiene

$$i\hbar \frac{d\psi}{dt} = E e^{-i(\omega t + \phi)} = H\psi. \quad (14.62)$$

avendo sfruttato l'equazione di Schrödinger che ci dice che $E\psi = H\psi$. Quest'ultima equazione prende il nome di **equazione di Schrödinger dipendente dal tempo** e rappresenta l'analogo di $\mathbf{F} = m\mathbf{a}$ della meccanica classica. Con quest'equazione, infatti, si può calcolare l'evoluzione temporale di uno stato, conoscendo lo stato iniziale.

Fisicast

La meccanica quantistica è argomento dei seguenti podcast di Fisicast:

Il transistor <http://www.radioscienza.it/2013/05/09/il-transistor/>

L'effetto fotoelettrico <http://www.radioscienza.it/2012/11/22/leffetto-fotoelettrico/>

La Meccanica Quantistica nel mio cellulare
<http://www.radioscienza.it/2013/03/22/la-meccanica-quantistica/>

Una storia esemplare

La storia della fisica delle particelle è esemplare da molti punti di vista: le tappe che hanno portato i fisici a formulare l'attuale modello della fisica delle particelle illustrano in maniera emblematica il progresso scientifico e le modalità con le quali si attua.

La nascita della fisica delle particelle si può far risalire ad anni diversi, che vanno dai primi del 1900 agli anni '30 del XX secolo, secondo le preferenze dei diversi autori. Noi stabiliremo la data di nascita della fisica delle particelle all'anno 1912, nel corso del quale il fisico austriaco **Viktor Hess** dimostrò l'esistenza dei **raggi cosmici**.

I raggi cosmici furono scoperti cercando di rispondere a una domanda all'apparenza del tutto irrilevante: perché gli oggetti elettricamente carichi, prima o poi si scaricano? Siamo certi che la maggior parte dei lettori penseranno che sia del tutto normale che ciò avvenga e che pochissimi avrebbero considerato l'opportunità di dare risposta a una domanda così apparentemente insignificante. E invece, come già accaduto in altre occasioni, il tentativo di rispondere a questa domanda diede vita a una serie di scoperte sorprendenti e alla nascita di una disciplina completamente nuova!

15.1 La scarica degli elettroscopi

All'inizio del XX secolo gli **elettroscopi** erano strumenti piuttosto diffusi nei laboratori di fisica. Già da tempo si era notato che, una volta caricati, dopo alcune ore perdevano la loro carica. Come sempre, quando si osserva un fenomeno, questo deve essere

spiegabile in termini fisici e la scarica di un elettroscopio è uno di questi. Se un elettroscopio si scarica vuol dire che perde progressivamente le cariche elettriche che si trovano distribuite sulla superficie delle sue parti conduttrici.

Per rimuovere tali cariche è necessaria la presenza di una qualche forza, di natura elettrica, che modifichi lo stato di carica del sistema in esame. Occorre dunque individuare la sorgente di tale forza. È abbastanza naturale aspettarsi che la sorgente debba essere una carica elettrica che attrae le cariche presenti sull'elettroscopio, strappandole da questo. Una tale carica si potrebbe naturalmente trovare all'interno dei laboratori dove si fanno gli esperimenti, per molti motivi. Vennero così avviate campagne di misura per determinare quali potessero essere le possibili sorgenti.

Gli elettroscopi venivano schermati con materiali diversi, portati in luoghi diversi, il più possibile lontano da cariche elettriche libere. In tutti i casi l'osservazione sperimentale era la stessa: tutti gli elettroscopi, indipendentemente dalle condizioni nelle quali si trovavano, andavano progressivamente scaricandosi.

La recente scoperta della **radioattività naturale** portò alcuni scienziati a ipotizzare che la scarica degli elettroscopi fosse da imputare alla presenza di materiali radioattivi, sempre presenti sulla crosta terrestre, che emettevano **raggi β** o **raggi α** che, essendo elettricamente carichi, avrebbero potuto interagire con gli elettroscopi provocandone la scarica. Si cercò allora di verificare quest'ipotesi, misurando il tasso di ionizzazione in luoghi nei quali l'abbondanza di elementi radioattivi era nota. Ci si aspettava, naturalmente, che la ionizzazione (e quindi la



Figura 15.1 Domenico Pacini al lavoro (1910).

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 15.2 Caricando elettrostaticamente un elettroscopio si osserva che, dopo un tempo più o meno lungo, perde tutta la sua carica [<http://www.youtube.com/watch?v=eI189swyr7Q>]. Anche se completamente isolato.

rapidità di scarica degli elettroscopi) fosse più intensa laddove gli elementi radioattivi erano abbondanti, come certe miniere. E invece il risultato fu inconcludente, perché non si riuscì a correlare la ionizzazione con l'abbondanza degli elementi ritenuti responsabili.

Da queste misure si poteva dedurre che la causa della ionizzazione non fosse da ricercarsi in elementi presenti nella crosta terrestre. Fu così che l'italiano **Domenico Pacini** decise di eseguire misure di ionizzazione sott'acqua, a diverse profondità. Pacini eseguì svariate misure sia in acqua salata, al largo di **Livorno**, che in acqua dolce, nel **Lago di Bracciano**, per escludere eventuali effetti dovuti alla presenza di sali disciolti nell'acqua. In entrambi i casi, Pacini ottenne il risultato secondo il quale la ionizzazio-



Figura 15.3 Viktor Hess a bordo del pallone aerostatico poco prima di partire per una delle sue campagne di misura.

ne diminuiva con la profondità. In un articolo [15] apparso sul **Nuovo Cimento** nel 1912, Pacini concludeva che “*esista nell'atmosfera una sensibile causa ionizzante, con radiazioni penetranti, indipendente dall'azione diretta delle sostanze radioattive del terreno.*”. I raggi ionizzanti, infatti, dovevano essere assorbiti dall'acqua e quindi, se fossero stati presenti al di sopra della superficie di questa, man mano che si scendeva in profondità si doveva misurarne sempre meno.

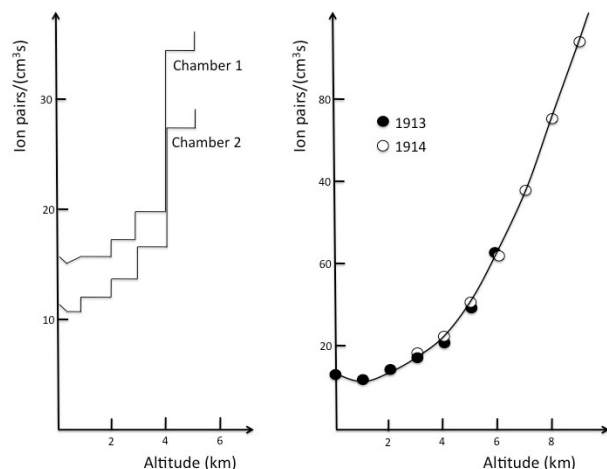


Figura 15.4 Misure di ionizzazione eseguite da Hess nel 1912, in funzione della quota raggiunta dal pallone con a bordo gli strumenti. A destra sono mostrate le misure, più precise, eseguite negli anni successivi da W. Kolhörster.

15.2 La scoperta dei raggi cosmici

Purtroppo per Pacini, la conclusione, pur corretta, non era completa. Egli, infatti, attribuì all'atmosfera il ruolo di *contenitore* degli agenti ionizzanti. L'austriaco **Viktor Hess**, invece, si spinse leggermente più in là, anche perché disponeva di una tecnologia più avanzata: i **palloni aerostatici**. Hess eseguì misure di ionizzazione a bordo di un pallone, a quote diverse, anche molto alte (fino ad alcuni km). La ionizzazione aumentava con la quota in modo evidente, come si vede bene dalla Figura 15.4, ottenuta riportando i dati dall'articolo originale [16] con il quale, sempre nel 1912, lo stesso Hess annunciò la scoperta.

Da queste misure Hess ricavò l'idea che la sorgente della radiazione ionizzante fosse esterna alla Terra. Secondo l'ipotesi di Hess la radiazione ionizzante aveva natura extra-terrestre: penetrava nell'atmosfera terrestre ed era parzialmente assorbita da questa. Per questa ragione, allontanandosi dal

limite esterno dell'atmosfera (dunque avvicinandosi al livello del mare) l'intensità di questa radiazione diminuiva. Doveva trattarsi di un tipo di radiazione molto più penetrante di quelle allora conosciute, che avrebbero dovuto essere completamente assenti al livello del mare. Non si trattava, dunque né di **radiazione α** , né di **radiazione β** .

Per questa radiazione di origine extra-terrestre fu coniato più tardi, da **Robert Millikan** il termine di **raggi cosmici**. Secondo il **Merriam-Webster Dictionary** la prima volta che tale termine comparve risale al 1925, quindi diversi anni più tardi rispetto alla loro scoperta. L'articolo in questione [17] è disponibile online sul **sito** dell'Accademia delle Scienze Statunitense.

15.3 Caratteristiche dei raggi cosmici

Dal momento che questi *raggi* provocavano effetti ionizzanti, dovevano essere costituiti di particelle elettricamente cariche. Ben presto si scoprì che le particelle in questione dovevano per lo più essere di carica positiva. Le misure, infatti, rivelarono che esisteva un'asimmetria nella direzione dalla quale i raggi cosmici apparivano provenire: in particolare sembrava che i raggi cosmici provenissero preferibilmente da ovest (**effetto est-ovest**). Fu il fisico **Bruno Rossi** a predire, nel 1930, questo fenomeno.

Com'è noto la Terra è circondata da un campo magnetico le cui linee di forza sono grosso modo parallele ai meridiani e sono dirette da sud a nord. I raggi cosmici arrivano da tutte le direzioni, per cui possiamo immaginare un flusso di particelle dirette in media verso la Terra. Usando la **regola della mano destra**, possiamo facilmente determinare la direzione e il verso della **Forza di Lorentz** agente su una particella di carica positiva. Disponendo le dita del palmo della mano in direzione del campo magnetico terrestre (quindi lungo un meridiano nella direzione sud-nord) e il pollice in direzione della velocità delle particelle, cioè verso la superficie terrestre, la forza di Lorentz è diretta perpendicolarmente al palmo della mano, perciò verso ovest. Se dunque i raggi

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 15.5 Raggi cosmici attraversano una [camera a scintilla](http://www.youtube.com/watch?v=HvEZbZc4XnA), che ne rende visibili le tracce [http://www.youtube.com/watch?v=HvEZbZc4XnA]. In una camera a scintilla una miscela di gas nobili (He e Ne) è ionizzata dal passaggio delle particelle cariche. Nel gas ci sono piani metallici tra i quali si applica una forte differenza di potenziale provocando una scintilla nel punto in cui il gas è stato ionizzato. Le scintille quindi si allineano lungo la traiettoria percorsa dalla particella che ha attraversato lo strumento.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 15.6 Spiegazione dell'effetto est-ovest nei raggi cosmici [http://www.youtube.com/watch?v=OZG6Bk45ZKs]. Questi appaiono provenire per lo più da ovest, diretti verso est. Il motivo è che i raggi cosmici primari sono protoni, quindi particelle positive, il cui moto è influenzato dalla presenza del campo magnetico terrestre.

cosmici sono prevalentemente carichi positivamente sembreranno provenire maggiormente da ovest. In effetti si osserva un flusso maggiore in direzione ovest-est.

Oggi sappiamo che la stragrande maggioranza dei cosiddetti **raggi cosmici primari**, quelli cioè che arrivano in prossimità della Terra dallo spazio, sono **protoni**. Questi costituiscono circa il 90–95 % del flusso totale di raggi cosmici. Il resto è quasi

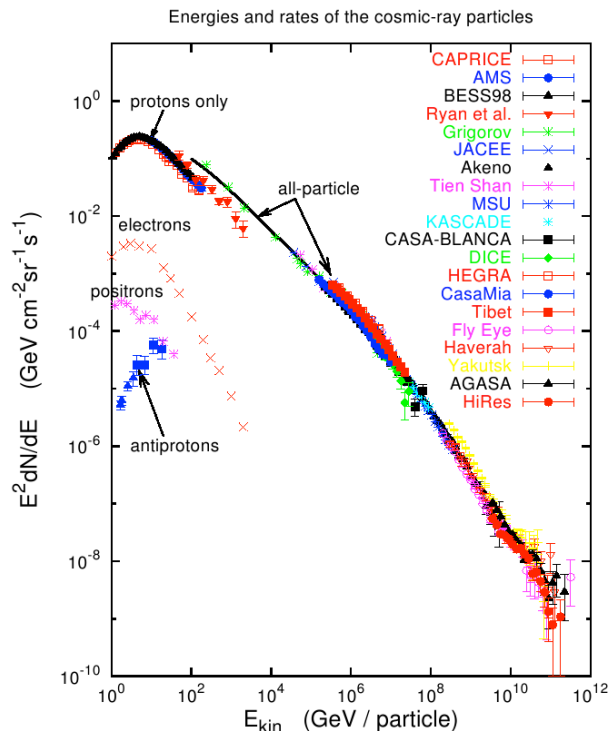


Figura 15.7 Spettro energetico dei raggi cosmici [18] come misurato da diversi esperimenti. Intorno ai 10^6 – 10^7 GeV si osserva un leggero cambiamento nella pendenza della curva (che è in scala logaritmica). Questa regione è convenzionalmente chiamata la regione del *gionocchio*: la forma dello spettro in questa scala, infatti, ricorda quella di una gamba leggermente flessa all'altezza del ginocchio.

completamente costituito di **particelle α** (nuclei di elio). Meno dell'1 % di essi è costituito di nuclei più pesanti e altrettanti sono gli **elettroni**.

Lo spettro energetico dei raggi cosmici (cioè il flusso per unità d'energia) è mostrato in Fig. 15.7. Si nota l'amplissimo intervallo di energie che arriva fino a 10^{12} GeV (corrispondenti a quasi 200 J!) per singola particella.

Si osserva anche un abbondante flusso di raggi γ

e X. I raggi cosmici primari collidono con i nuclei degli elementi che costituiscono l'atmosfera terrestre. Nell'urto si producono numerose nuove particelle chiamate **raggi cosmici secondari** che si dirigono verso la superficie (essenzialmente per la conservazione della quantità di moto). Parte di essi giunge al livello del mare dove il flusso di raggi cosmici secondari è pari a circa $100 \text{ m}^{-2}\text{s}^{-1}$.

I raggi cosmici provengono parzialmente dal Sole (si osserva un flusso maggiore in direzione di questo), ma la maggior parte di essi deve essere di origine extra-galattica, perché il flusso appare privo di direzionalità (in pratica non si osservano raggi cosmici provenire direttamente da una sorgente specifica, se non in misura relativamente modesta: i raggi cosmici appaiono provenire da tutte le direzioni) e le uniche potenziali sorgenti distribuite uniformemente attorno alla Terra sono le galassie. Le possibili sorgenti di raggi cosmici sono tutte quelle nelle quali si possono accelerare protoni e produrre fotoni. Le stelle sono una possibile sorgente, ma le energie in gioco nei processi termonucleari che avvengono al loro interno sono troppo basse per spiegare lo spettro osservato. Una volta prodotti nelle stelle, dunque, i raggi cosmici devono poter essere accelerati da qualche processo in grado di fornire loro le energie osservate. Le esplosioni delle *supernovae* potrebbero essere uno di questi processi.

La forma dello spettro suggerisce che i raggi cosmici siano accelerati, in media, in modo uniforme. Il primo modello di accelerazione fu pubblicato da **Enrico Fermi** nel 1949 [19].

Esercizio 15.1 *Il Modello di Fermi*

Una particella E viaggia per un tempo Δt in una regione di spazio dove può essere accelerata. Nel corso di questo viaggio ha una certa probabilità P di guadagnare energia ΔE . Il guadagno di energia è tanto più alto quanto più è lungo il tempo Δt di permanenza nella regione e tanto più alto quanto maggiore è la sua energia iniziale E , perciò $\Delta E = \alpha E \Delta t$ dove α è una costante di proporzionalità. Naturalmente, nello stesso tempo, la probabilità di perdere energia è $1 - P$. In questo caso, supponiamo che la particella esca dalla regione di spazio in cui può essere accelerata e la possiamo considerare persa (non giungerà mai nei dintorni della Terra).

Dimostra che lo spettro di energia atteso per le particelle sopravvissute è una legge di potenza. Per farlo scrivi l'energia guadagnata da una particella di energia E dopo aver attraversato una regione di spazio nella quale guadagna energia per un tempo lungo t . Quindi valuta la sua probabilità di sopravvivenza assumendo che, se la particella non guadagna energia avendo percorso un breve tratto, venga espulsa dalla regione accelerante. Il numero di particelle che giungono sulla Terra è proporzionale a questa probabilità (sarà il numero di particelle iniziali moltiplicato per la probabilità di sopravvivere nel viaggio dal punto in cui sono state prodotte, fino a noi).

Le particelle cariche possono essere accelerate se attraversano regioni in cui sono presenti campi magnetici variabili.

SOLUZIONE →

Chi l'ha ordinato?

La storia della fisica delle particelle, come spesso avviene, ha seguito un cammino talvolta tortuoso e costellato da una serie di errori di valutazione dei risultati sperimentali o da interpretazioni risultate poi non del tutto corrette. Non è negli scopi di questa pubblicazione ricostruire tale storia, perciò in questo capitolo si ripercorre una storia leggermente alterata rispetto a quanto realmente accaduto: una storia adattata al fine di meglio illustrare il processo che porta alla scoperta di un nuovo fenomeno e di far comprendere meglio la fisica che c'è dietro ogni scoperta.

16.1 Particelle penetranti

Negli anni successivi alla scoperta dei raggi cosmici si erano perfezionati alcuni strumenti per l'indagine scientifica, in grado di visualizzare la traccia prodotta da particelle cariche in moto.

Uno di questi strumenti era la **camera a nebbia** o **camera di Wilson**, inventata da **Charles Wilson**. La camera a nebbia è di fatto un recipiente contenente un gas **soprassaturo**, molto vicino al punto di condensazione. Quando una particella carica attraversa il vapore, questo tende a condensare proprio laddove la particella ha lasciato una traccia ionizzata, perché le particelle di gas ionizzate tendono ad attrarre elettrostaticamente altre particelle. Si forma quindi una serie di goccioline di liquido lungo la traiettoria seguita dalle particelle cariche che hanno attraversato lo strumento. Le goccioline, illuminate, hanno l'aspetto di una sottile nuvoletta bianca.

Nel 1929 il fisico russo **Dmitri Skobeltsyn** aveva ottenuto circa 600 fotografie di eventi in una camera a nebbia tenuta immersa in un campo magne-

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 16.1 Costruzione di una semplice camera a nebbia [<http://www.youtube.com/watch?v=qYhhmjYNwq4>]. Eseguite l'esperimento in una giornata secca e al buio, illuminando la camera di lato con una lampada da tavolo e ponendo un fondo nero dal lato opposto rispetto a quello da cui guardate. Dovrebbero bastare 100–150 ml di alcool. Se è troppo poco non vedrete nulla. Se è troppo, vedrete una sottile *pioggia* cadere in continuazione dalla spugna.

tico uniforme [20]. In queste fotografie si vedono le tracce di particelle cariche che percorrono traiettorie circolari per effetto della **Forza di Lorentz**. In 32 di queste, tuttavia, sono presenti tracce approssimativamente rettilinee, provenienti dall'esterno della camera. Le particelle che si osservano in queste foto sono cariche, avendo lasciato la traccia nello strumento (le particelle neutre non ionizzano il vapore e non provocano la condensazione). Dalla dimensione e densità delle goccioline se ne deduce che tali particelle devono avere una carica elettrica pari, in modulo, a quella dell'elettrone. Il fatto che vadano praticamente dritte vuol dire che possiedono un'enorme quantità di moto, dal momento che il **raggio di curvatura** r di una particella di carica q e massa m in un campo magnetico B è dato dalla formula

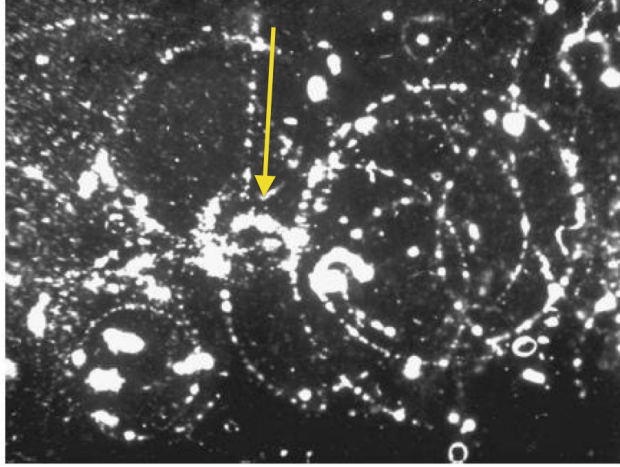


Figura 16.2 Una delle foto fatte da Skobeltsyn, in cui si vedono le tracce di particelle cariche che descrivono traiettorie circolari. Tra tutte le tracce se ne distingue una, accanto alla freccia gialla, del tutto simile alle altre, ma praticamente dritta, proveniente dall'esterno. È la prima fotografia della traccia di un raggio cosmico, probabilmente un *muone*.

$$r = \frac{\gamma m v}{q B} \quad (16.1)$$

dove γ è il **fattore relativistico di Lorentz**

$$\gamma = \frac{1}{\sqrt{1 - \left(\frac{v}{c}\right)^2}}. \quad (16.2)$$

La traccia appare dritta perché il raggio di curvatura è talmente ampio da non essere misurabile, per cui Skobeltsyn poté solamente stimare l'energia minima di queste particelle che doveva essere molto maggiore di quella delle particelle che stava studiando. Le particelle in questione provenivano principalmente da una direzione prossima allo zenit e fu immediato identificare queste tracce con quelle lasciate da particelle dei raggi cosmici.

16.2 L'ipotesi del neutrino

Nel 1930 **Wolfgang Pauli** propose di spiegare lo spettro energetico della **radiazione β** ipotizzando l'esistenza di una nuova particella [21]: il **neutrino**¹.

Il decadimento β consiste in un processo nel quale alcuni elementi, come il Cobalto 60 (^{60}Co), si trasformano spontaneamente in un elemento con la stessa massa, ma di specie atomica diversa: il Cobalto, ad esempio, si trasforma in Nichel 60 (^{60}Ni). Il Bismuto 210 (^{210}Bi) decade β tramutandosi in Polonio 210 (^{210}Po). In generale, partendo da un atomo ${}^A_Z N$, si produce un atomo con lo stesso peso atomico A , ma con diverso numero atomico Z , che differisce da quello originale ${}^A_Z N$ di un'unità:

$${}^A_Z N \rightarrow {}^A_{Z\pm 1} N + X \quad (16.3)$$

dove X rappresenta l'insieme degli altri prodotti della reazione. Il Nichel 60 ha lo stesso peso atomico del Cobalto 60, ma ha numero atomico 28, mentre il numero atomico del Cobalto è 27.

Nella trasformazione è emesso un elettrone (in passato gli elettroni erano noti con il nome di raggi β), in modo tale che la carica elettrica sia conservata. Infatti, il nucleo del Cobalto possiede 27 protoni, mentre quello del Nichel 28. Il Bismuto 210 ha lo stesso numero di neutroni e protoni del Polonio 210, ma il Polonio ha un protone al posto di un neutrone. La carica elettrica del nucleo, quindi, aumenta di un'unità. La carica elettrica è una grandezza conservata, pertanto non può cambiare nel corso del tempo. In effetti, quando il Cobalto si trasforma in Nichel, emette un elettrone, di carica elettrica opposta a quella del protone, in modo tale che, complessivamente, la carica elettrica dello stato iniziale e di quello finale sia la stessa.

Prima dell'ipotesi di Pauli si riteneva che il processo fosse del tipo

$${}^A_Z N \rightarrow {}^A_{Z+1} N + e^- \quad (16.4)$$

¹In realtà Pauli chiamava questa particella *neutrone*. Fu Fermi a ribattezzarla *neutrino* successivamente alla scoperta del neutrone.

Interazioni Deboli

Il decadimento β deve essere provocato da un'interazione di qualche tipo. La Fisica Classica insegna che le forze producono il cambiamento dello stato di un corpo che, sempre in Fisica Classica, è determinato quando se ne conoscano posizione e velocità. Per questa ragione applicando una forza a un corpo se ne cambia la velocità. Ma non sempre lo stato di un corpo è determinato da queste due grandezze. Non è così, ad esempio, in [Meccanica Quantistica](#), dove lo stato dipende da energia e momento angolare (che sono dunque le grandezze che cambiano applicando una forza). Nel caso dei decadimenti cambia anche la natura della particella: il cambiamento di stato consiste nel fatto che inizialmente abbiamo una particella di un tipo ferma, che si trasforma in altre particelle in moto. Il cambio di stato è mediato da una forza che non può essere la forza di gravità né quella elettromagnetica, che non possono modificare la natura delle particelle cui sono applicate.

All'interazione responsabile del decadimento si dà il nome di **interazione debole**, perché la sua intensità è molto minore di quella dell'interazione elettromagnetica: confrontando la probabilità che una particella interagisca per interazione elettromagnetica con quella che la stessa particella sia soggetta all'interazione debole si trova un rapporto pari a circa 10^{11} in favore delle forze elettromagnetiche.

e veniva interpretato come la trasformazione di un neutrone in un protone, con conseguente emissione di un elettrone. Se così fosse, però, l'energia con la quale l'elettrone viene emesso dovrebbe essere sempre la stessa. Infatti, supponendo che il neutrone sia fermo nel nucleo, nel momento in cui si trasforma in un protone e in un elettrone, questi si devono allontanare l'uno dall'altro in modo da conservare la quantità di moto (che inizialmente è nulla).

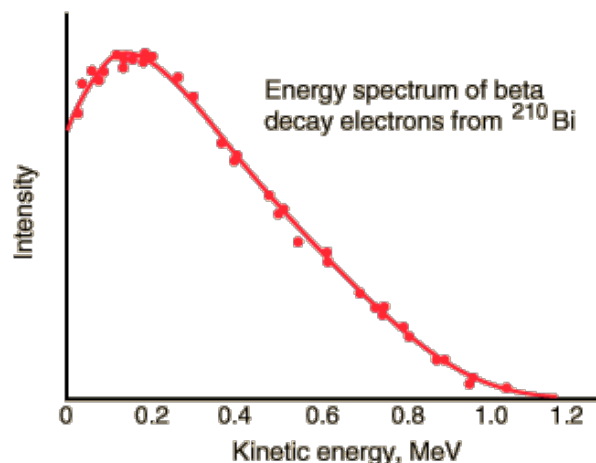


Figura 16.3 Spettro degli elettroni nel decadimento β del ^{210}Bi .

Esercizio 16.1 *Energia degli elettroni nel decadimento β .*

Sapendo che l'energia E di una particella di massa m e quantità di moto p , secondo la [relatività speciale](#), è pari a $E = \sqrt{p^2 c^2 + m^2 c^4}$, dimostra che la quantità di moto dell'elettrone proveniente da un decadimento β deve essere costante se il decadimento consiste nella trasmutazione di un neutrone in un protone e un elettrone ($n \rightarrow p + e^-$). Assumi che il neutrone sia inizialmente fermo, e imponi che l'energia e la quantità di moto siano conservate. Stima anche l'ordine di grandezza dell'energia massima che l'elettrone può assumere.

SOLUZIONE $\rightarrow$

Sperimentalmente si osserva tutt'altro: l'energia di un elettrone emesso in un decadimento β può assumere tutti i valori compresi tra 0 e un valore massimo che dipende dalle specie atomiche coinvolte, con uno spettro caratteristico (nella Figura 16.3 è mostrato lo spettro di un elettrone del decadimento del Bismuto 210), tanto che si pensò alla possibilità che l'energia non fosse conservata nelle interazioni deboli o alla scala subatomica.

Questa forma dello spettro suggerisce che il deca-

dimento avvenga almeno a tre corpi. In altre parole, la reazione deve essere del tipo

$${}^A_ZN \rightarrow {}^A_{Z+1}N + e^- + \nu \quad (16.5)$$

dove ν è una particella neutra e molto penetrante, di massa trascurabile. Se infatti l'energia disponibile E_0 è divisa tra tre particelle, ciascuna di esse può assumere un valore compreso tra 0 ed E_0 . La particella in questione deve essere di massa praticamente nulla, perché solo così si può giustificare il fatto che il valore massimo assunto dall'energia dell'elettrone coincide di fatto con l'energia calcolata assumendo che

$$E_0 \simeq (m_n - m_p - m_e)c^2. \quad (16.6)$$

Deve essere neutra perché non si riusciva a rivelare questa particella con nessuno dei rivelatori a ionizzazione disponibili, ma non poteva essere un fotone perché questi si potevano rivelare usando tecniche alternative. Doveva quindi trattarsi di una nuova particella: il **neutrino**.

16.3 L'antimateria

Nel 1933 **Carl Anderson** fece una scoperta [22] sorprendente! In una celebre fotografia in camera a nebbia (Fig. 16.4) si vede una particella (nell'immagine provenire dal basso) che percorre una traiettoria curva, perché la camera era tenuta in una regione nella quale era presente un campo magnetico (perpendicolare al piano della figura). La particella colpisce un assorbitore in piombo (la fascia scura orizzontale) di 6 mm di spessore e ne emerge avendo perso energia (si nota, infatti, che il raggio di curvatura è più stretto rispetto a quello della particella incidente). L'energia persa si può stimare dalla variazione del raggio di curvatura e risulta compatibile con quella persa da un elettrone. La curvatura assunta nel campo magnetico, però, era opposta a quella che ci si aspettava per una particella di carica negativa.

Fu subito chiaro che doveva trattarsi di una particella del tutto identica a un elettrone, con la stes-

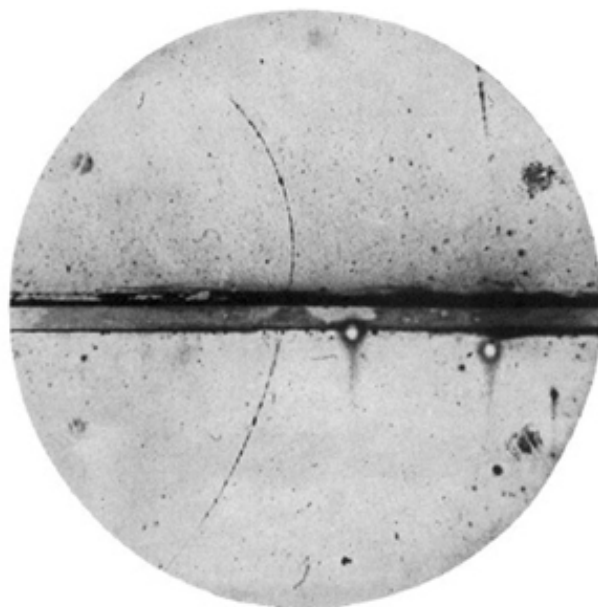


Figura 16.4 La prima fotografia di un positrone: una particella di antimateria.

sa massa, ma con carica elettrica opposta. A questa particella lo stesso Anderson diede il nome di **positrone**. Il positrone si indica col simbolo e^+ .

Si trattava della prima particella di **antimateria** scoperta. L'antimateria è costituita di particelle in tutto e per tutto uguali a quelle che compongono la materia ordinaria, ma con la carica elettrica opposta a quella di queste ultime. Per essere precisi, tutte le cariche, di qualunque natura, hanno segno opposto per materia e antimateria. Così, ad esempio, l'antiparticella del neutrino ν è l'antineutrino $\bar{\nu}$. Il neutrino non ha carica elettrica, ma possiede cariche di natura diversa (come il **numero leptonico**), il cui segno è opposto a quello della stessa carica del neutrino.

Le antiparticelle, in generale, si indicano con una barra sul simbolo della particella oppure indicando esplicitamente il segno della carica elettrica: e^+ nel caso del positrone oppure $\bar{p}$ nel caso dell'antiprotone.

L'esistenza di queste particelle di antimateria era stata **predetta** da **Paul Dirac** qualche anno prima, utilizzando semplici argomenti teorici. la scoperta

del positrone risultò dunque essere una brillante conferma della teoria.

Esercizio 16.2 Scoperta del positrone

Calcola il modulo e il verso del campo magnetico presente nell'esperimento di Anderson, sapendo che il positrone aveva un'energia di 63 MeV prima di attraversare il piombo e di 23 MeV dopo.

Usa la Figura 16.4 per ricavare la curvatura che le particelle assumono in campo magnetico, sapendo che lo spessore dell'assorbitore (la banda scura al centro) era di 6 mm.

SOLUZIONE →

16.4 La scoperta del muone

Negli anni successivi numerosi esperimenti avevano osservato tracce di particelle cariche che non erano ascrivibili né a elettroni, né a protoni. Le tracce sembravano essere prodotte da una particella che aveva una massa compresa tra quella dell'elettrone e quella del protone (si poteva desumere dall'intensità della ionizzazione) e carica uguale a $\pm e$ (indicando con e la carica del protone).

Una di queste osservazioni si deve a Paul Kunze [23] che nel 1932 aveva osservato alcune tracce sospette in camera a nebbia che aveva attribuito alle particelle che si osservavano nei raggi cosmici, già osservate da Blackett e Occhialini [24] poco tempo prima.

Anche Carl Anderson e Seth Neddermeyer avevano osservato eventi di questo tipo [25] e J. Street e E. Stevenson [26] ipotizzarono esplicitamente l'esistenza di una nuova particella: il muone (μ).

Pare che quando venne ufficialmente riconosciuta l'esistenza di questa nuova particella, nel corso di un congresso, il fisico Rabi esclamò "chi l'ha ordinato, questo?", come si fa al tavolo di un ristorante quando il cameriere, per errore, porta un piatto che nessuno ha chiesto. Nessuno, infatti, sentiva il bisogno di questa nuova particella per spiegare alcunché.

Il muone μ è una particella che si trova in due possibili stati di carica (μ^+ e μ^-) e ha una massa pari a circa 200 volte quella dell'elettrone. Analizzando alcuni eventi registrati su emulsioni nucleari si era scoperto che i muoni erano particelle instabili, con una vita media di circa $2 \mu\text{s}$, che decadevano lasciando una traccia carica attribuibile a un elettrone. Questa traccia aveva lunghezza variabile, il che indicava che l'elettrone emesso nel decadimento aveva quantità di moto variabile da evento a evento. Per le stesse ragioni per cui si fece l'ipotesi del neutrino, si poté stabilire che il muone decadeva secondo la reazione

$$\mu \rightarrow e + \nu + \bar{\nu} \quad (16.7)$$

cioè in un elettrone e due neutrini, che in seguito si scoprirono appartenere a due specie diverse: ν_e e $\bar{\nu}_e$. Naturalmente, per la conservazione della carica elettrica, il μ^+ decade in positroni e^+ e il μ^- in elettroni e^- .

16.5 La scoperta del pione

Pochi anni dopo la scoperta del muone, sempre nei raggi cosmici, venne scoperta (per primi da Lattes, Occhialini e Powell) una nuova particella, anch'essa in due stati diversi di carica. La nuova particella non poteva essere né un muone né un elettrone né un protone (né una delle rispettive antiparticelle) perché presentava caratteristiche diverse.

La ionizzazione prodotta dalle tracce di questa particella, ribattezzata **pione** o π , era molto simile a quella prodotta da un muone, quindi doveva avere grosso modo la stessa massa. In realtà era un pochino più pesante, perché negli esperimenti si osservavano i decadimenti dei pioni in muoni.

Le fotografie (Fig. 16.5) in emulsioni nucleari o in camera a nebbia mostravano tracce di pioni che si arrestavano in certi punti e da lì emettevano un muone (riconoscibile dal successivo decadimento e dalla ionizzazione). La lunghezza della traccia del muone era sempre la stessa, il che indicava che i muoni provenienti dal decadimento di un pione dovevano avere sempre la stessa energia. Per le ragioni

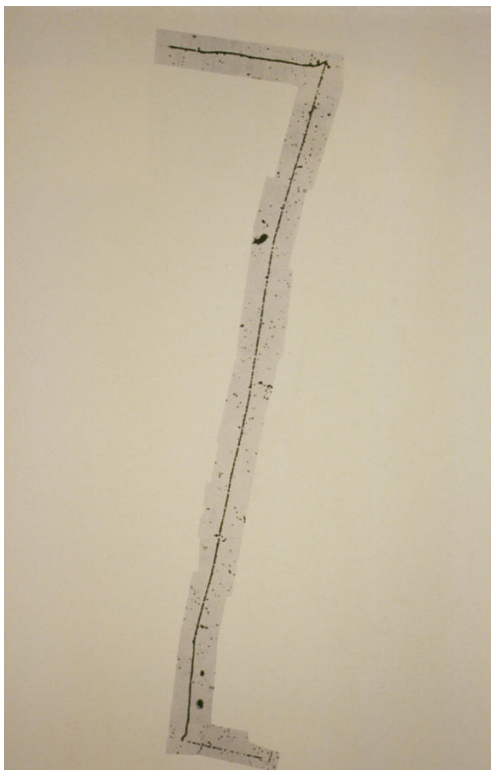


Figura 16.5 Il decadimento di un pione. Il π è quello in basso, che entra da destra, si ferma in un punto e da lì emette un muone (la traccia più lunga) che, a sua volta, decade nel punto in alto. Dai punti in cui avviene il decadimento del pione e del muone appare provenire una sola traccia, perché le altre particelle prodotte nella reazione sono neutrini che non lasciano tracce nello strumento.

esposte nel Paragrafo 16.2 questo significa che il decadimento del pione è un decadimento a due corpi (vedi anche l'Esercizio 16.1). Dalla parte opposta a quella in cui andava il muone non si rivelava mai nulla. Questo portò a concludere che la particella emessa insieme al muone doveva essere un neutrino, per cui il decadimento del pione è

$$\pi^{\pm} \rightarrow \mu^{\pm} + \nu_{\mu} \quad (16.8)$$

e ν_{μ} può essere sia un neutrino che un antineutrino

(secondo lo stato di carica del π).

Più tardi si scoprì che il pione esisteva anche in uno stato di carica neutro (π^0). Il π^0 decade in una coppia di fotoni:

$$\pi^0 \rightarrow \gamma\gamma \quad (16.9)$$

Il quadro delle particelle elementari si era complicato non poco. Con la scoperta di protoni, elettroni e neutroni si era creduto d'aver compreso la struttura della materia. La scoperta dei muoni e dei pioni, per non parlare di quella dei positroni, aveva reso tutto molto più difficile. Ma non era tutto...

16.6 La lambda e i mesoni K

Nel 1947 George Rochester e Clifford Butler intrapresero una lunga campagna di misurazioni eseguendo quasi 5000 fotografie di eventi in camera a nebbia per un totale di 1500 ore di osservazione. Analizzando quest'imponente quantità di dati trovarono un evento particolarmente interessante [27].

L'evento presentava due tracce, prodotte evidentemente da due particelle cariche, che formavano una V rovesciata, come quelle visibili nella Fig. 16.6.

Le tracce provenivano chiaramente dallo stesso punto e il vertice si trovava nel gas e non all'interno di un assorbitore di 3 cm di piombo piazzato nella camera. Questo fatto indicava che le due particelle non erano il prodotto di una collisione, ma di un decadimento, cioè della trasformazione di un'altra particella, presumibilmente prodotta dalla collisione di una particella nel piombo.

La presenza di un campo magnetico permetteva di misurare la quantità di moto e la carica elettrica delle tracce, che si dimostrarono avere carica opposta. La particella madre perciò doveva essere neutra. Dalla misura di quantità di moto si poté stabilire che la massa della particella madre doveva essere compresa tra 770 e 1600 volte la massa dell'elettrone. Nessuna delle particelle note fino ad allora aveva queste caratteristiche e si ritenne d'aver scoperto una nuova particella.

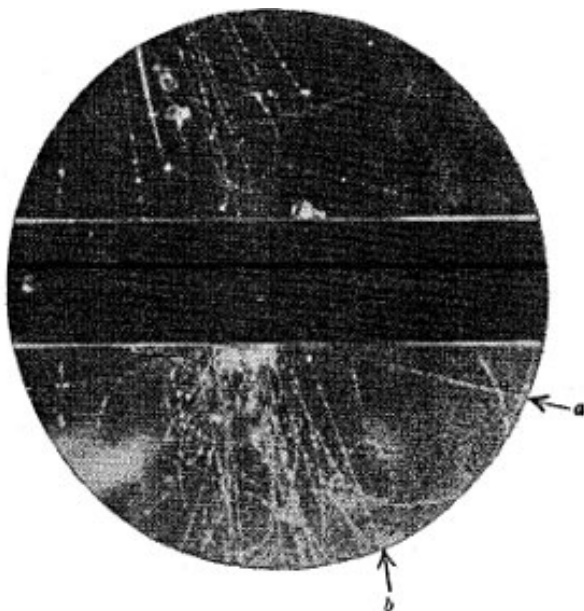


Figura 16.6 Uno degli eventi trovati da Rochester e Butler nei raggi cosmici. Sono evidenziate le tracce che si aprono a formare una V rovesciata.

Anche Hopper e Biswas scoprirono una particella dalle caratteristiche simili, nel 1950 [28]. Anche in questo caso si osservava (Fig. 16.7) una coppia di tracce che si aprivano a formare una V, ma la massa della particella neutra che si supponeva essere quella che dava origine all'evento era molto maggiore di quella della particella scoperta qualche anno prima da Rochester e Butler.

La particella scoperta da Hopper e Biswas fu chiamata Λ , proprio per la forma che assumevano le tracce negli eventi. Quella di Rochester e Butler fu battezzata K .

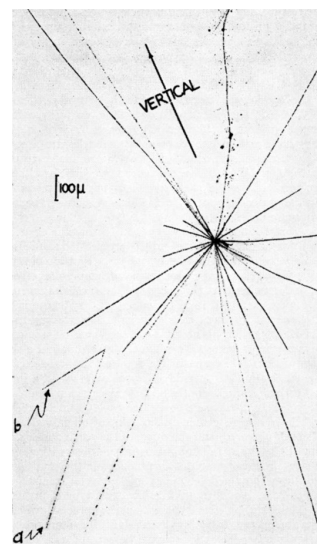


Figura 16.7 L'evento scoperto da Hopper e Biswas. Le tracce evidenziate dalle lettere a e b sono quelle prodotte dal decadimento di una particella Λ .

I nuovi numeri quantici

Con il proliferare delle particelle note, bisognava quanto meno capire perché i loro decadimenti erano quelli osservati. Quello che si poteva stabilire empiricamente era che in Natura esistevano regole di conservazione che impedivano il verificarsi di certi eventi che apparivano essere possibili sulla carta.

17.1 I leptoni

Non sembrerebbe esserci alcuna ragione per la quale il decadimento del muone dovrebbe procedere come un decadimento a tre corpi. Un decadimento del tipo $\mu \rightarrow e + \nu$ sembrerebbe del tutto plausibile. Evidentemente un tale decadimento deve essere vietato da qualche legge di conservazione. Si può supporre, in analogia a quanto accade con la conservazione della carica elettrica per cui il numero di cariche in unità della carica dell'elettrone resta costante, che esista una grandezza fisica conservata detta **numero leptonico**. Il muone e il suo neutrino devono avere lo stesso numero leptonico L_μ , mentre l'elettrone e il suo neutrino devono avere un numero leptonico diverso L_e .

Se il numero leptonico è conservato, in presenza di un μ^- nello stato iniziale il numero leptonico vale $L_\mu = +1$. Nello stato finale si deve avere lo stesso numero leptonico, per cui ci deve essere una particella che *porta* questo numero. Il neutrino muonico ν_μ evidentemente ha $L_\mu = +1$ e così un tale numero si conserva. Ma il neutrino non può essere l'unico prodotto della reazione: ce ne vuole almeno un altro. L'elettrone potrebbe essere uno di questi, ma e^- ha numero leptonico $L_e = +1$, dunque deve essere accompagnato da una particella che abbia $L_e = -1$

che è il suo antineutrino $\bar{\nu}_e$. Questo spiega perché si osservano le reazioni

$$\mu^- \rightarrow e^- + \bar{\nu}_e + \nu_\mu \quad (17.1)$$

e

$$\mu^+ \rightarrow e^+ + \nu_e + \bar{\nu}_\mu \quad (17.2)$$

mentre non si osserva il decadimento

$$\mu^- \rightarrow e^- + \bar{\nu}_e. \quad (17.3)$$

I muoni, gli elettroni e i rispettivi neutrini si dicono perciò appartenere alla classe dei **leptoni** (dal greco, che vuol dire *leggero*¹).

17.2 I barioni

In maniera del tutto analoga si può ritenere che esista un altro numero conservato detto **numero barionico**. Il neutrone e il protone hanno numero barionico pari a $B = 1$, per cui il neutrone può decadere solo se nello stato finale è presente un protone, per conservare questo numero. Poiché non può decadere in una sola particella (non si conserverebbero l'energia e la quantità di moto) nello stato finale deve essere presente almeno un'altra particella. Dal momento che la massa del neutrone è solo di poco più grande rispetto a quella del protone, il protone dello stato finale può essere accompagnato solo da particelle leggere come l'elettrone. Ma un elettrone ha un numero leptonico $L_e = +1$ e se ci fosse solo lui il numero leptonico non sarebbe conservato. Deve dunque esserci anche un antineutrino

¹Oggi questa parola ha perso il suo significato etimologico, esistendo anche leptoni *pesanti* come il τ .

che ha $L_e = -1$ per far sí che il numero leptonico complessivo sia nullo.

Ecco perché il decadimento è $n \rightarrow p + e^- + \bar{\nu}_e$. Il protone non decade perché è la particella più leggera a possedere un numero barionico. Anche i neutroni presenti nei nuclei atomici, solitamente, non decadono (altrimenti gli atomi, come li conosciamo, non si formerebbero). Il motivo è che i neutroni presenti nel nucleo *convertano* parte della loro massa in energia di legame con gli altri componenti del nucleo e non hanno più la massa sufficiente per produrre un protone e un elettrone. Il decadimento è proibito, dunque, dalla conservazione dell'energia.

Interazioni forti

Per spiegare perché i nuclei atomici sono stabili e non tendono a *rompersi* in mille pezzi è necessario ipotizzare l'esistenza di una forza molto più intensa di quella elettromagnetica. Due protoni nello stesso nucleo si respingerebbero infatti con una forza pari a circa

$$F_{em} = k \frac{q^2}{r^2} \simeq 9.0 \times 10^9 \left(\frac{1.6 \times 10^{-19}}{1 \times 10^{-15}} \right)^2 \simeq 230 \text{ N} \quad (17.4)$$

avendo usato per $r \simeq 10^{-15}$ m, il raggio tipico di un nucleo atomico. Perché i due protoni restino nel nucleo deve esistere una forza che li trattienga al suo interno più intensa di quella elettromagnetica, cui siano soggetti anche i neutroni (che non sono elettricamente carichi): la **forza o interazione forte**, appunto.

Le particelle che hanno un numero barionico sono particelle che risentono dell'**interazione forte**, che è quella responsabile del fatto che protoni e neutroni stanno insieme nel nucleo atomico (che altrimenti si disgregherebbe per effetto della repulsione elettrostatica). Per questa ragione neutroni e protoni sono **barioni**. Il termine barione deriva dal greco e vuol dire **pesante** (protoni e neutroni sono abbastanza più pesanti delle altre particelle di cui parliamo sopra), anche se oggi ha perso questo significato, es-

sendo state scoperte particelle non barioniche (cioè non simili ai protoni) più pesanti.

17.3 I mesoni

I pioni possono decadere in due corpi perché non hanno numero leptonico. Evidentemente non hanno neanche un numero barionico, perché nello stato finale non ci sono barioni. I pioni appartengono alla classe dei **mesoni** (come è facile immaginare, anche il termine mesone deriva dal greco e vuol dire “di medio peso”). Chiaramente non esiste un *numero mesonico*, perché nel decadimento del π non ci sono regole particolari che impediscono qualche decadimento altrimenti possibile per la conservazione dell'energia. Anche se appare strano che i pioni decadano in muoni invece che in elettroni (il decadimento $\pi \rightarrow e + \nu_e$ si osserva, ma con bassissima probabilità). In linea di principio il decadimento in elettroni dovrebbe essere favorito perché l'elettrone è più leggero ed è in un certo senso più facile produrlo².

17.4 Gli adroni

Barioni e mesoni, che risentono entrambi dell'**interazione forte**, appartengono a una classe più vasta detta degli **adroni** (*forte*, in greco). I barioni sono sempre particelle di **spin** semintero ($J = \frac{1}{2}$ o $J = \frac{3}{2}$), mentre i mesoni hanno sempre spin intero ($J = 0$ o $J = 1$).

In definitiva possiamo dividere le particelle in due classi, quella degli adroni e quella dei leptoni, sulla base del tipo d'interazione cui sono sensibili (i primi sia all'interazione forte che a quella **debole**, i secondi solo a quella debole). Gli adroni sono a loro volta divisi nella classe dei barioni (a spin semintero) e in quella dei mesoni (a spin intero).

²Questo comportamento anomalo si spiega con una caratteristica particolare dei neutrini: quella di esistere in un solo stato di spin. Non discutiamo questo fenomeno in questa sede.

Esercizio 17.1 *Caccia all'intruso*

Trova, tra le seguenti, le reazioni vietate dai principi di conservazione dei numeri quantici illustrati in questo capitolo, motivando la scelta.

$$\begin{array}{ll}
 p + p \rightarrow p + \bar{p} & p + p \rightarrow p + p + p + \bar{p} \\
 p + \pi^0 \rightarrow n + e^+ & p + n \rightarrow n + p + \pi^0 \\
 \pi^+ \rightarrow \pi^0 + \nu_e & \pi^- + p \rightarrow n + \pi^0 \\
 \mu^+ + n \rightarrow n + p + \bar{\nu}_\mu & \mu^- + e^- \rightarrow \bar{p} + e^- \\
 \pi^+ + \pi^- \rightarrow p + \bar{p} & \pi^+ + \pi^- \rightarrow p + n + \pi^-
 \end{array}$$

Controlla tutti i numeri quantici: carica elettrica, numero leptonico e numero barionico.

[SOLUZIONE →](#)

17.5 Classificazione in base allo spin

Le particelle, indipendentemente dalla loro natura, si possono anche classificare in base al loro **spin**, che è una grandezza fisica **quantizzata** il cui valore è sempre un multiplo intero di $\hbar/2$, dove $\hbar = h/2\pi$ è la **costante di Plank** h divisa per 2π . In molti casi, in Fisica delle Particelle, si usano le **unità naturali**, nelle quali $\hbar = 1$ ed è adimensionale. Pertanto i valori ammessi per lo spin, in queste unità, sono 0, $\frac{1}{2}$, 1, $\frac{3}{2}$, 2, e così via.

Le particelle a spin intero si chiamano **bosoni**. I pioni, ad esempio, sono bosoni. Tutti i mesoni sono bosoni. Ma non tutti i bosoni sono mesoni (ad esempio, il fotone, la particella che costituisce la luce) è un bosone, ma non è un mesone. I mesoni sono bosoni che subiscono l'**interazione forte**.

Le particelle a spin semintero, invece, si chiamano **fermioni**. L'elettrone, il muone, il protone e gli altri barioni sono fermioni. Da questo elenco si capisce che tutti i barioni sono fermioni, ma non tutti i fermioni sono barioni (ad esempio, il muone è un fermione, ma non è un barione).

Imitare la Natura

Lo studio dei raggi cosmici era diventato lo studio della materia che compone l'Universo. Per comprendere le proprietà delle nuove particelle occorreva eseguire numerose misure per ciascuna delle quali era richiesto un gran numero di **eventi**. Per raccogliere la statistica sufficiente a studiare un determinato tipo di particelle non si poteva far altro che attendere che la Natura le producesse e sperare nella fortuna. È l'*approccio del pescatore* che getta l'amo in mare e si mette in attesa che il pesce abbocchi. Se è fortunato prende almeno un pesce tra quelli desiderati, ma può darsi che non lo sia e all'amo potrebbe non abboccare nulla oppure qualche pesciottino di poco conto.

Una maniera più furba di procedere, perlomeno se si vuole essere certi di non saltare la cena, è quella di allevare il pesce, invece che di catturarlo. In altre parole, sarebbe molto più efficiente poter produrre le particelle desiderate in un laboratorio, scegliendo il tempo e la durata dell'esperimento e avendo la certezza di rivelare la maggior parte di quelle prodotte.

18.1 Gli acceleratori di particelle

Un modo per produrre le particelle consiste nell'imitare i processi che la Natura mette in atto per *riformarci* di raggi cosmici: occorre una qualche particella (i protoni dei raggi cosmici primari) da accelerare opportunamente, in modo tale da fargli raggiungere l'energia sufficiente affinché, collidendo con gli atomi di un bersaglio (nel caso dei raggi cosmici questi sono quelli dei gas che compongono l'atmosfera), producano le particelle desiderate. A valle del

bersaglio basterà piazzare dei rivelatori per osservarle. I vantaggi di questo approccio sono evidenti: si sa quando e dove le particelle sono prodotte, a quale angolo (o almeno in quale intervallo di angoli saranno abbondanti) e si può scegliere, in una qualche misura, la particella da produrre.

In effetti la produzione di nuove particelle dall'urto di una di queste con un nucleo può avvenire grazie alla [teoria della relatività](#), per la quale massa ed energia sono semplicemente due aspetti diversi della stessa grandezza fisica essendo che

$$E^2 = p^2 c^2 + m^2 c^4 \quad (18.1)$$

dove E è l'energia di una particella di massa m e quantità di moto p . Per particelle ferme $p = 0$ e si ritrova la celebre equazione

$$E = mc^2. \quad (18.2)$$

Disponendo di un'energia E si potrebbero **materializzare** particelle in numero tale che la somma delle loro masse non superi il valore $m \leq E/c^2$. In ogni caso, da un urto tra particelle di energia complessiva E possono venir fuori particelle la cui somma delle masse m e delle quantità di moto p sia tale da rispettare l'equazione (18.1).

Esercizio 18.1 *Energia di soglia*

Calcola l'energia minima necessaria affinché dall'urto di due protoni si possa produrre un pione neutro. La reazione da considerare è

$$p + p \rightarrow p + p + \pi^0. \quad (18.3)$$

Per farlo imponi la conservazione dell'energia e dell'impulso e osserva che la quantità $E^2 - p^2c^2$ è una costante indipendente dal sistema di riferimento usato per scrivere E e p . Puoi quindi imporre che questa differenza sia la stessa nel sistema del laboratorio in cui inizialmente uno dei protoni si muove e l'altro è fermo, e nel sistema di riferimento del centro di massa, in cui la somma delle quantità di moto è nulla.

SOLUZIONE →

I raggi cosmici primari sono costituiti per lo più di protoni. Possiamo procurarci protoni ionizzando idrogeno, il cui nucleo contiene un solo protone. Possiamo poi accelerarli usando un **acceleratore** facendogli acquisire energia sufficiente affinché, scontrandosi con i protoni fermi all'interno di un bersaglio¹, come fanno i raggi cosmici con i nuclei dei gas che compongono l'atmosfera, producano, per esempio, un π^0 attraverso la reazione

$$p + p \rightarrow p + p + \pi^0. \quad (18.4)$$

Il π^0 , una volta prodotto, decade in due fotoni. Ponendo dunque dei rivelatori di fotoni dopo il bersaglio dovremmo osservare l'arrivo in coincidenza di due fotoni ogni volta che si accende l'acceleratore. Misurando le proprietà di questi fotoni, quindi, possiamo risalire alle proprietà del π^0 che li ha generati e che ci interessano.

¹un materiale molto usato è il **berillio** per le sue proprietà termiche e meccaniche

Esercizio 18.2 *La massa invariante*

Supponi di conoscere la quantità di moto $\vec{p}_i$ e l'energia E_i di due fotoni ($i = 1, 2$). Il fotone è una particella a massa nulla $m_\gamma = 0$. Supponendo che i due fotoni siano il risultato del decadimento di una particella, calcola la massa che deve avere questa particella per produrre i due fotoni nello stato dato. Per farlo osserva che la differenza $E^2 - p^2c^2$ è pari a m^2c^4 , dove m è la massa di una particella, se E e p sono energia e quantità di moto di questa particella. Imponi la conservazione dell'energia e della quantità di moto, poi calcola la differenza $E^2 - p^2c^2$ nello stato iniziale e in quello finale.

SOLUZIONE →

È così che si dà inizio a una nuova campagna d'esperimenti, eseguiti stavolta in laboratori attrezzati con acceleratori di protoni, usati per *sparare* particelle su bersagli e produrre così le particelle d'interesse. La fisica si sposta così all'interno dei laboratori dove si costruiscono acceleratori sempre più sofisticati e potenti.

In certi casi è possibile produrre particelle come i π^+ e i π^- per urto tra protoni, da accelerare successivamente per usarle come proiettili per studiare reazioni del tipo

$$\pi^\pm + p \rightarrow X \quad (18.5)$$

e

$$\pi^\pm + n \rightarrow X \quad (18.6)$$

dove X rappresenta uno dei possibili stati finali (che dipende dall'energia dei pioni, e deve avere **numero leptonico** complessivo nullo e **numero barionico** uguale a 1).

L'adozione di queste tecniche permetterà lo studio intensivo della fisica delle particelle e consentirà nuove scoperte, come illustrato nei capitoli seguenti.

Studiare le particelle

Avendo a disposizione un fascio di particelle possiamo studiare sostanzialmente due aspetti: quanto intensamente le particelle interagiscono con la materia che le circonda (quindi con quale probabilità le particelle presenti nel fascio cambiano direzione, energia o natura attraversando un blocco di materiale), oppure quanto rapidamente le particelle prodotte si trasformano spontaneamente (**decadono**) in altre particelle.

Ogni volta che si definisce una grandezza, in fisica, se ne deve dare la **definizione operativa**: si deve cioè dire come in pratica si esegue la loro misura. In questo capitolo spieghiamo come si definiscono le grandezze fisiche caratteristiche per descrivere quanto sopra e come si procede operativamente per assegnare loro un valore.

19.1 Sezione d'urto

Nell'urto tra una particella e un bersaglio si possono misurare diverse grandezze fisiche. Una di queste, molto semplice da misurare, almeno in linea di principio, è il numero di particelle che hanno interagito col bersaglio.

Il processo d'urto ha carattere probabilistico, per cui alcune delle particelle del fascio attraverseranno il bersaglio senza subire alcun effetto, mentre altre saranno deviate oppure spariranno per lasciare il posto a nuove particelle.

Sia N il numero di particelle inviate sul bersaglio. La variazione nel numero di particelle $\Delta N = N' - N$, dove N' è il numero di particelle che hanno attraversato il bersaglio senza subire assorbimento, deviazioni o perdita di energia, è chiaramente proporzionale al numero di particelle incidenti:

$\Delta N \propto -N$ (il segno meno indica che si ha una diminuzione nel numero di particelle nel fascio).

Questo numero sarà, in modulo, tanto maggiore, quanto più spesso è il bersaglio, perché più il bersaglio è spesso, più le particelle incidenti possono interagire con esso. Quindi possiamo scrivere

$$\Delta N \propto -N\Delta x \quad (19.1)$$

dove Δx rappresenta lo spessore del bersaglio. È anche chiaro che maggiore è la densità del bersaglio ρ , maggiore sarà la probabilità d'interagire con i suoi nuclei, pertanto abbiamo che

$$\Delta N = -\sigma N\rho\Delta x \quad (19.2)$$

dove σ è una costante di proporzionalità che chiamiamo **sezione d'urto**. Osserviamo che le **dimensioni fisiche** di ρ sono quelle di un volume alla meno uno (ρ rappresenta il numero di particelle nel bersaglio per unità di volume), mentre Δx ha le dimensioni fisiche di una lunghezza. N e ΔN sono adimensionali e perciò σ deve avere le dimensioni di una superficie: $[\sigma] = [L^2]$. Si misura, quindi, in m^2 nel **SI**. Per comodità si definisce l'unità di misura della sezione d'urto come il **barn**, che corrisponde a 10^{-24} cm^2 e si indica col simbolo **b**.

La sezione d'urto dipende dal tipo di processo e dalla specie della particella interagente e può dipendere dall'energia E della particella. Per misurare σ si prende un fascio di N particelle e lo si invia su un bersaglio di spessore noto Δx , di cui sia nota la composizione e quindi la densità ρ . Misurando il numero N' di particelle *diffuse* dal bersaglio, si misura la differenza $\Delta N = N' - N$ e si calcola

$$\sigma = -\frac{\Delta N}{N\rho\Delta x} \quad (19.3)$$

(che è un numero positivo perché ΔN è negativo). Oltre che una sezione d'urto per l'assorbimento di particelle, come in questo caso, possiamo definire una sezione d'urto di produzione. Se nell'urto tra una particella del fascio e quella del bersaglio si produce una nuova particella, contando il numero N_p di queste particelle possiamo definire la sezione d'urto di produzione come

$$\sigma = \frac{N_p}{N\rho\Delta x}. \quad (19.4)$$

La produzione di nuove particelle, in effetti, sarà tanto più probabile quanto maggiore è lo spessore Δx del bersaglio e la sua densità ρ . Inoltre, più particelle invio sul bersaglio e più ne produco di nuove.

La sezione d'urto misurata si può quindi confrontare con i modelli teorici delle interazioni. Usando un modello estremamente semplice possiamo pensare al processo di diffusione delle particelle di un fascio, da parte degli atomi del bersaglio, come all'urto di palline rigide. In questo caso la sezione d'urto ha un'interpretazione molto semplice:

$$\sigma \simeq \pi R^2 \quad (19.5)$$

dove R^2 rappresenta il raggio della sfera che rappresenta la particella bersaglio. Si tratta, naturalmente, di un raggio efficace, cioè di qualcosa che ha le dimensioni fisiche di una lunghezza, ma che non possiamo interpretare direttamente come il *raggio della particella*, ma piuttosto come il raggio entro il quale le interazioni che danno luogo al processo si fanno sentire.

Quando una particella proiettile attraversa un bersaglio si possono manifestare diverse interazioni: quelle di natura elettromagnetica, che ben conosciamo, e quelle dovute alle forze **forti** e **deboli**.

Le forze deboli, come dice il nome, sono molto poco intense e possiamo trascurarle per i nostri scopi. Le forze forti, al contrario, sono parecchio intense e quindi sono quelle che fanno sentire di più i

loro effetti. Affinché una particella subisca gli effetti della forza forte deve possedere una **carica forte** e quindi non deve essere un **leptone**. I protoni o i pioni subiscono la forza forte, pertanto le loro interazioni saranno dominate da questa (possiamo perciò trascurare gli effetti dell'interazione elettromagnetica).

Per protoni e pioni si misurano sezioni d'urto dell'ordine di $10^{-26} \text{ cm}^2 = 0.01 \text{ b}$.

Nel caso in cui si facciano interagire elettroni con la materia, le forze forti sono assenti, quindi possiamo assumere che questi interagiscono solo per interazione elettromagnetica. Le sezioni d'urto tipiche sono dell'ordine di $10^{-31} \text{ cm}^2 = 10^{-7} \text{ b} = 0.1 \mu\text{b}$.

Come ci aspettavamo, la sezione d'urto per interazione forte è molto maggiore di quella per interazione elettromagnetica.

Dall'equazione 19.2 si ricava che l'intensità del fascio di particelle diminuisce esponenzialmente con lo spessore attraversato

$$N(x) = N(0)e^{-\sigma\rho x} \quad (19.6)$$

dove $N(x)$ rappresenta il numero di particelle del fascio che non hanno ancora interagito alla profondità x e $N(0)$ quelle iniziali. Osservando che le dimensioni fisiche $[\sigma\rho]$ del prodotto di σ per ρ sono quelle di $[L^2L^{-3}] = [L^{-1}]$ possiamo definire

$$\sigma\rho = \frac{1}{\lambda} \quad (19.7)$$

dove λ è qualcosa che ha le dimensioni fisiche di una lunghezza $[\lambda] = [L]$ che si chiama **lunghezza d'interazione** e scrivere che

$$N(x) = N(0)e^{-\frac{x}{\lambda}}. \quad (19.8)$$

19.2 Vita media

Quando si ha a che fare con particelle instabili, una grandezza interessante da misurare è la **vita media**, che caratterizza il **decadimento**. Il decadimento di una particella consiste nella sua trasformazione in due o più particelle. Dal momento che in ogni processo fisico si conservano energia e quantità di moto,

il decadimento in una particella non può avvenire in un *canale* in cui sia presente una sola particella. In altre parole non esiste il decadimento $a \rightarrow b$ in cui la particella a si trasforma spontaneamente nella particella b . Se infatti la particella a è ferma, l'energia dello stato iniziale è $E = m_a c^2$, dove m_a è la massa di a e la quantità di moto è nulla. Se $m_b \neq m_a$ l'energia non è più conservata, a meno che la differenza di energia $(m_a - m_b)c^2$ non vada in energia cinetica per b , ma in questo caso la quantità di moto $p \neq 0$. D'altra parte, se $m_b = m_a$, $a = b$ e non c'è stato alcun decadimento.

Una particella dunque può decadere solo se esistono almeno due particelle la cui somma delle masse sia inferiore alla massa della particella madre. Un caso interessante si ha nel decadimento $a \rightarrow b + b$, quando le due particelle figlie hanno la stessa massa m_b . In questo caso la somma delle energie cinetiche di queste ultime deve essere pari a $(m_a - 2m_b)c^2$. Poiché però anche la quantità di moto si conserva, $\vec{p}_a + \vec{p}_b = 0$ e quindi le due particelle dello stato finale hanno la stessa velocità in modulo, ma sono emesse in direzioni opposte.

Dato un fascio con N particelle iniziali, trascorso un tempo Δt , le particelle ΔN decadute (che hanno cioè cambiato natura) sono proporzionali a quelle inizialmente presenti e al tempo trascorso, per l'appunto, quindi

$$\Delta N = -\alpha N \Delta t. \quad (19.9)$$

Il coefficiente di proporzionalità α deve avere le dimensioni fisiche di un tempo alla meno uno $[\alpha] = [T^{-1}]$ perché sia N che ΔN sono adimensionali. Per rendere evidente questo fatto conviene definire

$$\alpha = \frac{1}{\tau} \quad (19.10)$$

dove τ è una grandezza fisica che ha le dimensioni di un tempo e che è proprio quella che si chiama vita media.

Il processo di decadimento è un processo **stocastico**, che avviene in modo casuale secondo le leggi della statistica. Non possiamo predire quando e in quali particelle decadrà una data particella, ma

soltanto i possibili **canali** di decadimento, le percentuali relative (**branching ratio** o **rapporto di diramazione**) e il tempo medio nel corso del quale è possibile osservare un decadimento.

Di sicuro devono essere rispettate tutte le leggi di conservazione. Così abbiamo visto che i protoni non possono decadere in neutroni perché non hanno la massa sufficiente, e i muoni non possono decadere in due corpi ($\mu \rightarrow e + \nu_e$) perché questo decadimento non conserva il **numero leptonico**.

L'equazione (19.9) che definisce la vita media ha la stessa forma dell'equazione (19.2) e quindi, anche l'equazione che ci dice come varia N in funzione del tempo t **deve avere** la stessa forma di quella che dice come N varia in funzione di x :

$$N(t) = N(0)e^{-\frac{t}{\tau}}. \quad (19.11)$$

Quest'equazione ci dice che particelle prodotte simultaneamente non decadono tutte nello stesso momento. Circa $1/3$ di esse ($1/e$) decade in un tempo pari a τ , mentre in un tempo pari a 2τ ne saranno decadute $1 - \exp(-2)$ e cioè l'86 %. Più tempo si attende meno particelle sopravvissute si trovano. Per misurare τ quindi si prendono $N(0)$ particelle e si attende un tempo t abbastanza lungo. Si contano i decadimenti $N_d(t)$ avvenuti in questo tempo e si ricava $N(t) = N(0) - N_d(t)$. È facile vedere che

$$\log \frac{N(t)}{N(0)} = -\frac{t}{\tau} \quad (19.12)$$

e dunque

$$\tau = -t \frac{1}{\log \frac{N(t)}{N(0)}}. \quad (19.13)$$

Naturalmente, trattandosi di un processo statistico, si misura questa grandezza per diversi tempi t più volte e se ne calcola la media.

Esercizio 19.1 *Il decadimento del muone*

Per misurare il tempo di decadimento del muone **Marcello Conversi** e **Oreste Piccioni** [29] nel 1944 usarono tre rivelatori posti uno sull'altro. Tra quello più in alto e quello di mezzo si poneva un assorbitore in piombo. Quando i segnali emessi da questi rivelatori scattavano in coincidenza, si poteva essere certi che un muone aveva attraversato lo strumento (gli elettroni o i pioni¹ non riuscirebbero ad attraversare il piombo). Lo strato di rivelatori più basso era separato da quello di mezzo da un assorbitore in ferro. Si misuravano le coincidenze ritardate tra i rivelatori più in alto e questo. Se un muone si fermava nel ferro e successivamente decadeva, il numero di coincidenze ritardate doveva seguire la legge del decadimento esponenziale. Le coincidenze misurate nel corso di questo esperimento erano le seguenti (abbiamo riportato solo parte dei dati sperimentali):

Ritardo (μs)	Frequenza (Hz)
0.00	3.47 ± 0.4
1.00	2.42 ± 0.3
1.97	1.26 ± 0.2
3.80	0.86 ± 0.17

Calcola la vita media del muone.

SOLUZIONE →

Questa grandezza fisica si può misurare per tutte le particelle instabili, come i muoni, i pioni, e i neutroni. I π^0 decadono in due fotoni con una vita media dell'ordine dei 10^{-16} s. I pioni carichi, invece, decadono in tempi molto più lunghi, in muoni e neutrini: 10^{-8} s. I muoni, a loro volta, decadono in elettroni e neutrini con una vita media dell'ordine di un paio di μs e il neutrone ha una vita media di quasi 900 s. Si vede subito che l'intervallo di valori è molto ampio, tuttavia possiamo fare alcune considerazioni che ci portano a concludere che cer-

¹All'epoca dell'esperimento, in ogni caso, il pione non era stato ancora scoperto.

ti decadimenti siano dovuti all'**interazione debole**, mentre altri a quella elettromagnetica.

Lo spazio delle fasi

Lo **spazio delle fasi** è uno spazio nel quale si possono rappresentare tutte le possibili configurazioni di un sistema fisico e la sua *ampiezza* è una misura del numero di possibili configurazioni dati certi vincoli. In fisica delle particelle lo spazio delle fasi misura il numero di configurazioni che due o più particelle possono assumere a parità di condizioni. Nel decadimento di una particella in due corpi, ad esempio, lo spazio delle fasi è limitato dal fatto che le due particelle figlie devono avere quantità di moto uguale e opposta ed energia pari alla massa della particella madre (in unità naturali). Così in un decadimento di una particella di massa M in due particelle di massa m_1 e m_2 , l'energia della particella 1 è data da

$$E_1 = \frac{M^2 - (m_2^2 - m_1^2)^2}{2M}. \quad (19.14)$$

Nel caso particolare $m_2 = m_1$, E_1 è fissata e vale $M/2$: lo spazio delle fasi si riduce a un punto. Nel caso in cui esistano più canali di decadimento ci possono essere più valori per E_1 , che dipendono dalla differenza di massa tra le particelle figlie e lo spazio delle fasi si allarga. Maggiore è la differenza di massa, più ampio è lo spazio delle fasi.

Nel caso di un decadimento a tre corpi ciascuna particella può trasportare una frazione della quantità di moto, purché la somma vettoriale delle tre quantità di moto sia nulla. Quindi esiste un intervallo di valori permessi per la particella 1 e per la particella 2, mentre la particella 3 è obbligata ad assumere il valore determinato dalla conservazione della quantità di moto. Lo spazio delle fasi dunque è maggiore perché esistono più configurazioni rispetto al caso precedente.

Sicuramente il decadimento del neutrone è mediato dall'interazione debole ed è ragionevole aspettarsi che tutti i decadimenti che coinvolgano neutrini

siano mediati dalla stessa forza (i neutrini non sono carichi e sono *leptoni* quindi *sentono* solo l'effetto della forza debole). Quindi anche il decadimento dei muoni e dei pioni deve essere attribuibile a questa interazione. L'enorme differenza tra le vite medie è presumibilmente dovuta a fattori che dipendono dall'ampiezza dello *spazio delle fasi*. In effetti, il neutrone e il *protone* hanno masse molto simili, quindi esistono solo pochissime configurazioni energetiche permesse (l'energia a disposizione di elettrone e neutrino è pari alla differenza di massa tra le particelle in questione moltiplicata per la velocità della luce). Invece la differenza di massa tra i muoni e gli elettroni è molto maggiore, quindi è relativamente *facile* produrre configurazioni permesse dalla conservazione dell'energia (l'intervallo di energie che possono assumere i due neutrini è ampio). La differenza tra il tempo di vita medio del muone e di quello del pione non è così grande e si può facilmente attribuire al fatto che il decadimento del pione avviene in due corpi, per cui è più facile realizzare la configurazione giusta.

Il decadimento del π^0 , invece, non coinvolge neutrini e quindi è attribuibile all'interazione elettromagnetica, dal momento che nello stato finale ci sono solo fotoni, che rappresentano proprio la radiazione elettromagnetica (la luce, che è un'*onda elettromagnetica*, si può anche interpretare come un flusso di fotoni). In questo caso l'interazione è più intensa e il tempo di vita medio più corto.

Le risonanze

Usando **acceleratori di particelle** si possono così produrre altre particelle, come i pioni, inviando protoni accelerati su un bersaglio. I pioni carichi così prodotti possono essere a loro volta accelerati (purché lo si faccia prima che decadano) e inviati su un altro bersaglio. Si può così studiare, ad esempio, la sezione d'urto del processo $\pi + p$ o del processo $\pi + n$ (il bersaglio è certamente costituito di protoni e neutroni).

20.1 Urti tra particelle

A valle del bersaglio si pongono dei rivelatori con i quali si può, ad esempio, misurare la sezione d'urto di produzione di nuove particelle in funzione dell'energia del fascio incidente o dell'angolo di diffusione.

Esercizio 20.1 *Il decadimento del pione*

Calcola quanto spazio hai a disposizione per costruire un acceleratore in grado di raccogliere e accelerare i pioni prodotti nell'urto di protoni su un bersaglio, prima che questi decadano.

SOLUZIONE →

Se, ad esempio, si mandano dei π su un bersaglio, qualora i π non interagissero col bersaglio, questi seguirebbero una traiettoria rettilinea. Se quindi si misura la sezione d'urto di produzione di pioni a diversi angoli θ rispetto alla direzione di volo delle particelle del fascio, si avrebbe

$$\sigma(\theta) = 0 \quad (20.1)$$

per ogni $\theta \neq 0$, dal momento che tutte le particelle del fascio si muoverebbero in avanti senza essere deviate. In realtà i pioni interagiscono con i nuclei del bersaglio, quindi la sezione d'urto che si misura è diversa da zero. In altre parole, mettendo un contatore a un angolo θ rispetto alla direzione del fascio si misura un certo numero N_p di *nuove* particelle. Poco importa se si tratta, in realtà, di particelle già presenti nel fascio, che hanno semplicemente cambiato direzione in seguito all'urto con un nucleo. Non possiamo seguire le particelle individualmente, perciò non ha molto senso chiedersi se si tratti di una particella del fascio deviata o se nell'urto la particella presente nel fascio sia andata distrutta e ne sia stata prodotta un'altra, della stessa natura, con caratteristiche cinematiche diverse. Per noi, *nuova* significa una particella che prima dell'urto non c'era nello stato cinematico nella quale la osserviamo.

È abbastanza intuitivo capire che il numero di queste nuove particelle diminuisce all'aumentare dell'angolo di diffusione. Le diffusioni a piccolo angolo, in effetti, sono più probabili rispetto a quelle a grande angolo (non fosse altro che perché la quantità di moto è conservata e le velocità dei prodotti finali si devono sommare in maniera tale da conservare la quantità di moto in avanti). In ogni caso uno si aspetta che l'andamento di $\sigma(\theta)$ in funzione di θ sia caratteristico della specie di particella proiettile e di bersaglio. In linea di principio la sezione d'urto σ potrebbe dipendere anche dall'energia E della particella incidente ($\sigma = \sigma(\theta, E)$) ma non ci si aspettano variazioni brusche della sezione d'urto al variare dell'energia.

Quello che invece si osserva è un andamento della sezione d'urto in funzione dell'energia della particella incidente che varia lentamente con l'energia,

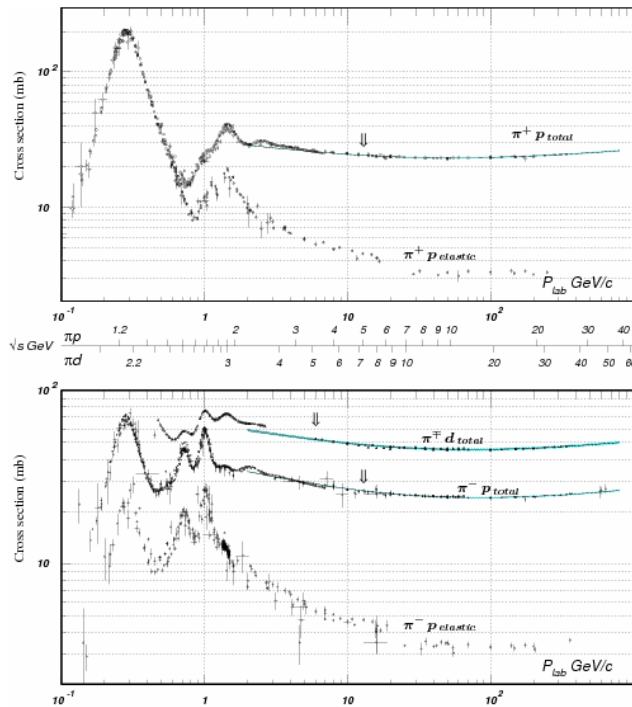


Figura 20.1 Sezione d'urto $\pi + p$ in funzione della quantità di moto del pione.

ma presenta, a certi valori caratteristici dell'energia, evidenti picchi nei quali la sezione d'urto può aumentare, rispetto al valore di base, anche di un ordine di grandezza o più. Nella Figura 20.1 si vede l'andamento della sezione d'urto del processo $\pi^+ p$ per diversi valori della quantità di moto p_{lab} del pione nel sistema di riferimento del laboratorio.

Sono riportate le sezioni d'urto elastica (quella del processo $\pi + p \rightarrow \pi + p$) e totale (quella del processo $\pi + p \rightarrow X$ dove X rappresenta un qualunque stato finale).

Appaiono evidenti una serie di valori di p_{lab} in corrispondenza dei quali la sezione d'urto aumenta parecchio. In certi casi, ad esempio per $p_{lab} \simeq 0.3$ GeV/c, la sezione d'urto elastica $\pi^+ + p$ passa da un valore dell'ordine di qualche mb a valori dell'ordine di 200 mb, mentre quella totale passa da 20 a 200 mb. Anche nel caso della reazione $\pi^- + p$ si osservano alcuni picchi, tutti in corrispondenza degli stessi valori di p_{lab} .

Questi picchi, detti **risonanze**, s'interpretano co-

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 20.2 L'osservazione di una nuova particella la cui vita media non è sufficiente a rivelarla direttamente si può fare attraverso la ricostruzione di quella che si chiama la sua *massa invariante*. Ricostruendo la massa dell'ipotetica particella che, decadendo, ha dato origine a due (o più) nuove particelle osservate nel rivelatore, si osserva un picco nella distribuzione di questa grandezza. [<https://www.youtube.com/watch?v=K2LfcwYzvBE>].

me la produzione di nuove particelle che successivamente decadono. Il motivo è semplice. Consideriamo l'urto elastico

$$\pi^+ + p \rightarrow \pi^+ + p \quad (20.2)$$

dove le due particelle nello stato finale si trovano in uno stato cinematico diverso da quello iniziale (nel senso che hanno una quantità di moto diversa). Se il processo che dà origine a questa reazione è un *banale* urto ci aspettiamo che la maggior parte delle particelle finisca in avanti, a un angolo piccolo rispetto alla direzione di volo del π^+ iniziale. Comunque a un angolo $\theta = \theta_0$ vedremo arrivare particelle con una certa frequenza, che diminuisce all'aumentare di θ . Se però nell'urto si produce una particella che successivamente decade in $\pi^+ + p$ i prodotti del decadimento non avranno più memoria della direzione iniziale del π^+ . I prodotti di decadimento potranno essere emessi ad angoli qualunque, in particolare ad angoli grandi. Quindi vedremo un forte aumento del numero di particelle osservate in corrispondenza dell'angolo θ_0 , quando l'energia E_π del π^+ incidente è tale da consentire la produzione di una particella con massa pari all'energia E_{cm} nel centro di massa

divisa per la velocità della luce al quadrato.

Esercizio 20.2 *Produzione di risonanze*

Calcola l'energia che deve avere un fascio di π^+ per produrre una risonanza di massa M . In particolare calcola la massa che potrebbe avere una risonanza inviando pioni con quantità di moto $p = 300$ MeV su protoni fermi.

In corrispondenza di questo valore si vede un picco nella sezione d'urto. La massa della risonanza in questione è di 1232 MeV.

SOLUZIONE →

Le particelle così prodotte devono essere altamente instabili. Infatti i prodotti di decadimento appaiono provenire direttamente dal punto d'interazione. La risonanze hanno quindi una vita media brevissima che si può stimare essere dell'ordine di 10^{-24} s.

La loro produzione è chiaramente mediata dalle **interazioni forti**. I pioni infatti devono interagire con i protoni, che sono particelle che subiscono l'interazione forte. D'altra parte anche i muoni interagiscono con i protoni, ma in questo caso le probabilità d'interazione sono molto più basse perché i muoni sono leptoni che non risentono dell'interazione forte. Anche il decadimento delle risonanze deve essere provocato dall'interazione forte, perché i tempi sono molto più brevi rispetto a quelli dell'interazione elettromagnetica e dell'**interazione debole**.

Esercizio 20.3 *Produzione di Δ*

Indica attraverso quali reazioni si possono produrre le quattro particelle Δ con fasci di pioni carichi.

SOLUZIONE →

Tra le risonanze si possono annoverare la Δ^- , la Δ^0 , la Δ^+ e la Δ^{++} , con carica elettrica, rispettivamente, pari a e , zero, $-e$ e $-2e$, dove e è la carica dell'elettrone.

20.2 La massa invariante

Secondo la **teoria della relatività** l'energia E di una particella di massa m e quantità di moto p è data da

$$E^2 = p^2 c^2 + m^2 c^4. \quad (20.3)$$

Nel caso in cui abbiamo a che fare con N particelle di massa m_i e quantità di moto p_i , $i = 1, \dots, N$, l'energia complessiva del sistema è la somma delle energie $E = E_1 + E_2 + \dots + E_N$ e la quantità di moto totale dalla somma di queste $\vec{p} = \vec{p}_1 + \vec{p}_2 + \dots + \vec{p}_N$. In questo caso possiamo scrivere che

$$E^2 = p^2 c^2 + M^2 c^4 \quad (20.4)$$

dove p è il modulo della quantità di moto totale e M è qualcosa che ha le dimensioni fisiche di una massa. Se consideriamo un insieme di particelle ferme è evidente che la somma delle loro energie $E = (m_1 + m_2 + \dots + m_N) c^2 = M c^2$. Quindi, almeno in questo caso, $M = m_1 + m_2 + \dots + m_N$.

L'energia E e la quantità di moto p di una particella sono grandezze fisiche che dipendono dal sistema di riferimento nel quale sono calcolate. Ad esempio: nel sistema di riferimento solidale con una particella, la sua energia vale $E = E_0 = m c^2$ e la sua quantità di moto $p = 0$. Se però la particella è in moto rispetto all'osservatore la sua energia $E > E_0$ e $p > 0$. La sua massa, però, non può dipendere dal sistema di riferimento scelto per misurarla. Si dice che la massa è un **invariante relativistico**. La differenza

$$E^2 - p^2 c^2 = m^2 c^4 \quad (20.5)$$

dunque, è un invariante: assume cioè sempre lo stesso valore in ogni sistema di riferimento. Se dunque abbiamo una particella ferma di massa M che decade in N particelle, ciascuna con energia E_i e quantità di moto $\vec{p}_i$, poiché M è invariante deve essere

$$E^2 - p^2 c^2 = M^2 c^4 \quad (20.6)$$

La radice di questa differenza¹ si chiama **massa invariante** ed è utile per determinare la massa di una particella che è decaduta in altre particelle di cui si conoscano energia e quantità di moto. Infatti, la massa invariante del sistema di particelle figlie deve coincidere con la massa della particella madre.

Quando dall'urto di un pione con un protone si produce una risonanza Δ , questa decade subito dopo sempre in una coppia pione-protone. Misurando la quantità di moto delle due particelle figlie p_π e p_p se ne può ricavare l'energia $E_\pi = \sqrt{p_\pi^2 c^2 + m_\pi^2 c^4}$ e $E_p = \sqrt{p_p^2 c^2 + m_p^2 c^4}$. Se si calcola la quantità

$$(E_\pi + E_p)^2 - (\vec{p}_\pi + \vec{p}_p)^2 c^2 \quad (20.7)$$

si ottiene proprio $m_\Delta^2 c^4$: il quadrato della massa della particella madre (moltiplicata per c^2). Per convincersene basta calcolare esplicitamente le grandezze in gioco. Nel sistema di riferimento in cui la Δ è ferma $\vec{p}_\pi = -\vec{p}_p$ e quindi $(\vec{p}_\pi + \vec{p}_p) = 0$. La massa invariante al quadrato quindi è data semplicemente da

$$M^2 c^4 = E_\pi^2 + E_p^2 + 2E_\pi E_p \quad (20.8)$$

Per la conservazione dell'energia deve anche essere che

$$E_\pi + E_p = m_\Delta c^2 \quad (20.9)$$

e facendo il quadrato di quest'ultima equazione si vede subito che $M = m_\Delta$.

Dato un sistema di N particelle prodotte dal decadimento di un'altra particella, conoscendone la cinematica, possiamo sempre calcolarne la massa invariante. Ora supponiamo di aver osservato in un esperimento diversi eventi in cui si trovano due particelle (che per semplicità supponiamo identiche). Queste due particelle possono essere il risultato del decadimento di una particella di massa M oppure possono essere state prodotte in un altro modo (per esempio nello stesso evento potrebbero essersi

sovrapposte due reazioni, ciascuna delle quali produce la particella in questione). Nel caso in cui siano le figlie di una particella di massa M la loro massa invariante sarà pari a Mc^2 , altrimenti sarà un numero a caso compreso tra 0 e E dove E è l'energia complessiva nello stato iniziale.

Dal momento che c è una costante, le masse delle particelle si possono indicare, invece che in unità di massa, in unità di energia. Un protone, ad esempio, ha una massa di circa 1.67×10^{-27} kg, che, moltiplicata per $c^2 = (3 \times 10^8)^2 = 9 \times 10^{16} \text{ m}^2\text{s}^{-2}$ dà

$$m_p c^2 \simeq 1.67 \times 10^{-27} \times 9 \times 10^{16} \simeq 15 \times 10^{-11} \text{ J} \quad (20.10)$$

che trasformato in elettronvolt (ricordiamo che $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$) vale

$$m_p c^2 \simeq 15 \times 10^{-11} \text{ J} \frac{1 \text{ eV}}{1.6 \times 10^{-19} \text{ J}} \simeq 9.4 \times 10^8 \text{ eV}. \quad (20.11)$$

In definitiva la massa del protone si può esprimere come circa 10^9 eV ($\simeq 1 \text{ GeV}$). In fisica delle particelle si usa esprimere le masse in unità di energia. Per ottenere le masse in kg basta dividere il valore in unità di energia per c^2 .

Se si hanno a disposizione numerosi eventi di questo tipo si può perciò calcolare la massa invariante per ogni coppia di particelle (nel caso di decadimenti a due corpi; la massa invariante si può calcolare anche per tre, quattro o più particelle nello stato finale) e costruire un **istogramma** delle frequenze con cui si presentano certi valori di massa invariante. Si ottiene un grafico in cui sull'asse delle ascisse sono riportati i possibili valori di massa invariante (in unità di energia) e in ordinata il numero di volte in cui si è trovato un determinato valore. Molti degli eventi saranno casuali: le due particelle trovate e considerate potenziali figlie di una particella più pesante non hanno in realtà alcuna relazione l'una con l'altra, quindi daranno origine a masse invarianti casuali comprese tra 0 e l'energia massima disponibile, ma in altri casi la massa invariante assumerà proprio il valore corrispondente alla particella che,

¹A rigore si dovrebbe chiamare "massa invariante" questa radice divisa per c^2 , ma dal momento che le due quantità differiscono solo per una costante moltiplicativa possiamo usare lo stesso nome per M e per Mc^2 .

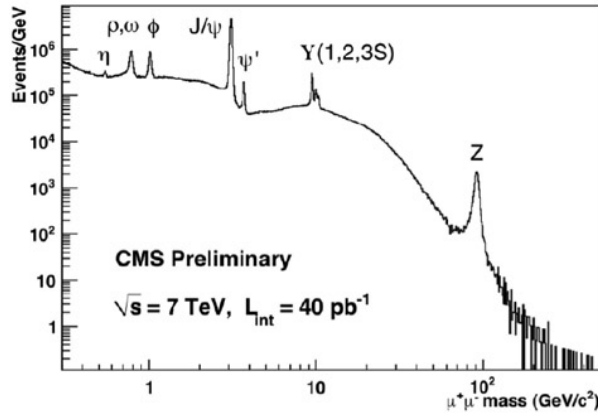


Figura 20.3 Distribuzione della massa invariante calcolata per una coppia di muoni $\mu^+\mu^-$ osservati all'esperimento **CMS** a **LHC**. Sul fondo piatto si osservano diversi picchi corrispondenti alla massa di altrettante particelle che decadono in una coppia di muoni. Notate la scala orizzontale, nella quale le masse sono espresse in GeV/c^2 . Il fondo in genere è decrescente con l'energia perché la probabilità di osservare coppie casuali di alta massa invariante è minore rispetto a quella di osservare coppie casuali di bassa massa invariante.

delle coppie $\pi + p$ che emergono dall'urto si trovano dei picchi in corrispondenza delle loro masse. Per questa ragione anche Λ e K si possono considerare risonanze.

decadendo, ha dato origine a quelle dello stato finale. Il grafico (Fig. 20.3) perciò avrà l'aspetto di un *fondo* continuo sul quale si stagliano alcuni **picchi** nella posizione corrispondente alla massa delle particelle madri.

Per estensione i picchi corrispondenti si chiamano anch'essi **risonanze**. In generale la scoperta di un nuovo picco su un fondo piatto di masse invarianti corrisponde alla scoperta di una nuova particella la cui massa coincide con la posizione del picco.

Nel grafico della sezione d'urto della reazione $\pi + p$ non si osservano picchi in corrispondenza della massa di particelle come la Λ o la K (il motivo è spiegato nel Cap. 21), tuttavia calcolando la massa invariante

Le particelle strane

Con gli acceleratori di particelle si possono studiare con un certo dettaglio le particelle scoperte nei raggi cosmici. In particolare le Λ e i K si possono produrre in numero praticamente arbitrario in laboratorio e questo permette di misurare la sezione d'urto di produzione e i tempi di decadimento.

Questo è strano, perché le leggi della fisica sono invarianti per inversioni temporali.

21.1 I decadimenti della Λ

Facendo queste misure si scopre che la sezione d'urto per la produzione di Λ e di K è quella tipica delle risonanze, perciò queste particelle sono prodotte per [interazione forte](#) attraverso le reazioni

$$\pi^- + p \rightarrow \Lambda + X \quad \pi^- + p \rightarrow K + X. \quad (21.1)$$

dove X rappresenta altre possibili particelle nello stato finale. I decadimenti di queste particelle, invece, hanno tempi tipici delle [interazioni deboli](#). La Λ decade principalmente nel canale

$$\Lambda \rightarrow \pi^- + p \quad (21.2)$$

per interazione debole e questo risultò subito *sospetto*, **strano**. Naturalmente non c'è niente di strano nel fatto che una particella sia prodotta per interazione forte e decada per interazione debole: i pioni, ad esempio, sono abbondantemente prodotti nelle interazioni forti e decadono per interazione debole. Anche i neutroni sono palesemente particelle che si possono produrre per interazione forte, ma decadono per effetto dell'interazione debole. Ma nessuna di queste particelle decade attraverso un'interazione diversa da quella che ne determina la produzione nello stesso canale in cui avviene quest'ultima.

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 21.1 Il moto di un pendolo è identico in avanti e all'indietro, perché la *caduta* del grave dipende dalla stessa interazione: la gravità. Se interviene un'interazione diversa a uno dei due estremi della traiettoria, il moto non è più simmetrico [<http://www.youtube.com/watch?v=IF2f4KYgfUc>].

Esercizio 21.1 *La produzione dei K*

Con un acceleratore di particelle possiamo scegliere l'energia da dare a un fascio di pioni da inviare su un bersaglio, in modo tale da garantire la produzione delle particelle che si desiderano studiare.

Calcola l'energia minima che dovrebbe avere un pione per produrre, attraverso l'urto con un protone fermo, una particella K . Nei dati sperimentali non si vedono risonanze a questa energia. Il motivo è illustrato nel Paragrafo 21.2.

SOLUZIONE →

Se si guarda un film alla TV, si capisce subito se il film è riprodotto normalmente o se le immagini vanno all'indietro nel tempo. Ma questo avviene solo perché le scene del film coinvolgono sistemi complessi composti di un numero sterminato di particelle elementari. Se riprendessimo con una telecamera l'*oscillazione di un pendolo* non potremmo stabilire se il filmato è riprodotto in avanti o all'indietro. In particolare, nel caso del pendolo, il tempo impiegato dalla massa appesa al filo per tornare indietro rispetto al punto d'inversione del moto è lo stesso ai due estremi della traiettoria, perché l'interazione che provoca il moto (la gravità) agisce nello stesso modo nei due punti. Questo tempo sarebbe diverso se a uno degli estremi intervenisse un'interazione

diversa (in particolare se in uno degli estremi fosse impedito alla gravità di fare il suo dovere).

Allo stesso modo, se possiamo produrre una Λ per interazione forte (quindi con alta probabilità) nell'urto tra un pione e un protone, il suo decadimento in un pione e un protone dovrebbe avvenire in tempi brevi con una probabilità analoga, dal momento che, invertendo la direzione del tempo, non dovremmo poter distinguere tra produzione e decadimento. Invece così non è. Per questo la Λ venne definita una particella **strana**.

L'unico modo di spiegare questo comportamento strano è di ammettere, come nel caso del pendolo, che nel decadimento l'interazione forte, responsabile della produzione, non possa fare il suo lavoro, lasciando il compito di far decadere la particella a un'altra interazione: quella **debole**.

Ma come mai l'interazione forte non può far decadere la Λ ? Evidentemente deve esistere una qualche grandezza fisica che è conservata nelle interazioni forti e non lo è nelle interazioni deboli. Potremmo pensare a una qualche caratteristica della particella che le interazioni forti *vedono*, e che, al contrario, per le interazioni deboli è irrilevante. Un po' come la carica elettrica, che è una caratteristica che determina il comportamento delle interazioni elettromagnetiche, ma che è del tutto irrilevante per le interazioni gravitazionali. Per la gravità un elettrone e un positrone sono identici, mentre per le interazioni elettromagnetiche no. Il fatto è che la gravità non provoca il cambiamento della natura delle particelle. Se lo provocasse, un elettrone potrebbe tranquillamente trasformarsi in un positrone.

Deve dunque esistere una specie di **carica** conservata nelle interazioni forti, che le interazioni deboli in un certo senso non *vedono*. Questa carica venne chiamata **stranezza**. Si dice che la Λ possiede una carica di stranezza (convenzionalmente pari a -1), che deve essere conservata nelle interazioni forti, ma può non esserlo nelle interazioni deboli. Perciò la Λ non può decadere secondo la reazione

$$\Lambda \rightarrow \pi^- + p \quad (21.3)$$

attraverso l'interazione forte perché né il pione né

il protone possiedono una carica di stranezza, che dunque non sarebbe conservata. Poiché però questa carica è irrilevante per le interazioni deboli, la Λ può decadere in questo modo attraverso la mediazione di quest'ultima interazione. Affinché si possa produrre una Λ per interazione forte, tuttavia, è necessario che si conservi la stranezza. Perciò sarebbe altrettanto impossibile osservare la reazione

$$\pi^- + p \rightarrow \Lambda \quad (21.4)$$

perché neanche in questo caso si conserverebbe la stranezza (che nello stato iniziale è nulla, mentre nello stato finale vale $S = -1$).

21.2 Produzione associata

La produzione potrebbe avvenire solo se nello stato finale fossero presenti almeno due particelle con stranezza opposta. In effetti, studiando meglio la produzione di particelle strane, si trova che le Λ sono sempre prodotte in associazione ai K . La reazione che si osserva è sempre del tipo

$$\pi^- + p \rightarrow \Lambda + K \quad (21.5)$$

seguita dal successivo decadimento delle Λ e dei K per interazione debole. I K , quindi, devono possedere una stranezza pari a $S = +1$. In questo modo, lo stato finale ha stranezza complessiva $S = 0$ e la reazione è possibile conservando la stranezza. Il decadimento di entrambe le particelle non può avvenire per interazione forte, che conserva la stranezza, ma può avvenire per interazione debole. I K , in effetti, decadono con tempi tipici delle interazioni deboli, in due o tre pioni:

$$K \rightarrow \pi^+ + \pi^- \quad K \rightarrow \pi^+ + \pi^- + \pi^0. \quad (21.6)$$

Studiando le reazioni agli acceleratori si scoprirono molte altre particelle strane, alcune delle quali con stranezza pari a multipli interi di quella della Λ . Ad esempio, la particella Ω^- ha addirittura stranezza $S = -3$. Questa particella si può produrre per interazione forte, ad esempio, solo in associazione a tre

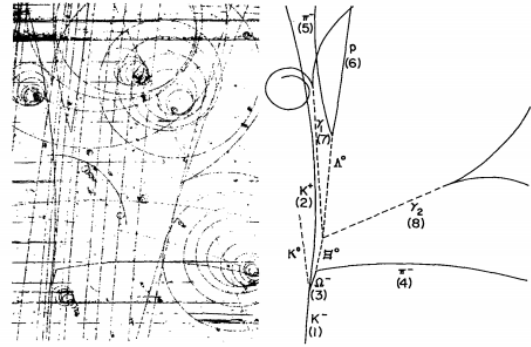


Figura 21.2 Produzione e decadimento di un Ω^- . Un K^- con stranezza $S = -1$ (in basso) urta un protone, producendo un Ω^- , un K^+ e un K^0 , conservando la stranezza. Il barione Ω decade quindi in $\Xi^0 + \pi^-$. Il barione Ξ^0 decade quindi in una Λ e un π^0 che subito si trasforma in due fotoni γ , che convertono in una coppia e^+e^- . L'ultimo decadimento nella catena è quello della Λ che decade in un protone e un pione. Nota la curvatura che le tracce delle particelle cariche assumono nel campo magnetico dell'esperimento.

particelle con stranezza $+1$ (o a una con stranezza $+2$ e una con stranezza $+1$). La Ω^- decade per interazione debole attraverso una complessa catena che conduce alla produzione, nello stato finale, di tre pioni e un protone. Inizialmente si ha il decadimento

$$\Omega^- \rightarrow \Xi^0 + \pi^- . \quad (21.7)$$

La Ξ^0 , un'atra particella strana, neutra con stranezza $S = -2$, decade poi in una Λ accompagnata da un pione neutro

$$\Xi^0 \rightarrow \Lambda + \pi^0 \quad (21.8)$$

e infine la Λ decade secondo la solita reazione $\Lambda \rightarrow \pi^- + p$.

La produzione associata spiega perché non si osservano risonanze nelle reazioni $\pi + n$ e $\pi + p$ quando l'energia del pione è sufficiente a produrre una

Λ o un K . Le reazioni citate, infatti, non possono produrre una Λ o un K singoli, ma devono per forza produrre queste due particelle in associazione, insieme.

II Modello a Quark

Attraverso lo studio intensivo delle possibili reazioni tra particelle e i loro decadimenti, furono scoperte moltissime nuove particelle. Il quadro si era ulteriormente complicato rispetto a quello, semplicissimo, in cui protoni, neutroni ed elettroni erano le uniche particelle elementari necessarie per spiegare la composizione della materia nell'Universo.

Oltre a queste tre particelle e ai pioni, si erano scoperte le quattro Δ prive di stranezza, tre particelle con stranezza $S = -1$ (Σ^- , Σ^0 e Σ^+), la Ξ nei due stati di carica Ξ^- e Ξ^0 , con stranezza $S = -2$, la Ω^- , con $S = -3$, la Λ e i K con stranezza rispettivamente $S = -1$ e $S = +1$.

Si scoprirono, inoltre, tre particelle simili alle Σ , ma più pesanti, che vennero chiamate Σ^* : Σ^{*-} , Σ^{*0} e Σ^{*+} . E due particelle simili alle Ξ , chiamate Ξ^* , negli stati di carica Ξ^{*-} e Ξ^{*0} .

Oltre a queste particelle si erano trovati due K carichi (K^+ e K^-), e una particella neutra chiamata η simile al π^0 .

Si era anche scoperto che esistevano K neutri con stranezza $S = +1$ e con stranezza opposta $S = -1$ per cui uno dei due doveva essere l'antiparticella dell'altro: K^0 e $\bar{K}^0$.

Per la conservazione del numero barionico le Λ dovevano essere barioni, mentre i K dovevano essere mesoni, come i pioni e la η . Le Σ e le Ξ sono barioni, così come le loro copie più pesanti Σ^* e Ξ^* , la Ω^- e le Δ .

Un quadro così complesso sembrava inspiegabile, fino a quando Murray Gell-Mann e Susumu Okubo proposero di considerare queste particelle come a loro volta composte di particelle più piccole, chiamate quark.

Tre quark per Muster Mark!

Il nome *quark* dato ai componenti elementari di alcune particelle è stato coniato da Murray Gell-Mann, il padre del Modello a quark.

Il termine è stato preso in prestito da Gell-Mann dal testo di *Finnegans Wake* di James Joyce, nel quale figura il brano

Three quarks for Muster Mark!

Sure he hasn't got much of a bark

And sure any he has it's all beside the mark.

Sembra che a Gell-Mann piacesse il suono di questa parola (quark) nel poema di Joyce, e il fatto che si facesse riferimento a tre quark, sembrò a Gell-Mann un buon motivo per scegliere questo nome, giacché servivano proprio tre quark per spiegare lo spettro delle particelle osservate.

La maggior parte dei fisici pronuncia la parola *quark* come *quork*, come faceva Gell-Mann (anche se probabilmente Joyce l'avrebbe pronunciata *quark*, per far rima con *Mark* e *bark*).

22.1 Tre nuove Tavole Periodiche

Che gli elementi chimici non fossero particelle elementari, ma composti di nuclei positivi ed elettroni, era stato suggerito dal fatto che le proprietà chimiche degli elementi consentivano di disporli nella Tavola Periodica di Mendeleev.

Allo stesso modo si potevano disporre tutte le nuove particelle scoperte su opportune Tavole,

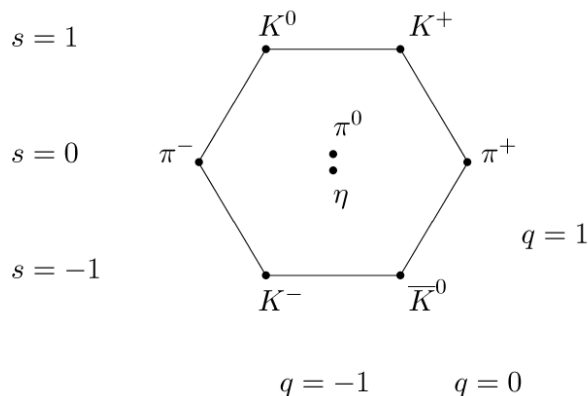


Figura 22.1 L'ottetto di mesoni.

scegliendo due numeri quantici quali indici della Tavola: la carica elettrica e la stranezza.

I mesoni, ad esempio, si potevano disporre come nella Figura 22.1, a formare il cosiddetto **ottetto di mesoni**. L'ottetto è una Tavola nella quale trovavano posto gli otto mesoni fino ad allora scoperti: i quattro K , i tre pioni e l' η . Nella prima riga comparivano quelli con stranezza $S = +1$, nella seconda quelli con stranezza $S = 0$ e nella terza quelli con stranezza $S = -1$. Le particelle poi si dispongono nello schema ai vertici di un esagono in modo tale da avere la stessa carica elettrica lungo linee oblique parallele a una delle diagonali. Così il K^- e il π^- sono allineati lungo uno dei lati dell'esagono, K^0 , η , π^0 e $\bar{K}^0$ lungo la diagonale parallela a questo lato e K^+ e π^+ lungo il lato opposto.

In maniera del tutto analoga, otto tra i **barioni** conosciuti, si potevano disporre in un altro ottetto usando gli stessi numeri quantici, come mostrato in Fig. 22.2.

Come nel caso dei mesoni, in cui π^0 e η occupano la stessa posizione nella Tavola, nel caso dei barioni Λ e Σ^0 possiedono gli stessi numeri quantici secondo i quali la Tavola è organizzata. Le particelle in questione differiscono solo per la massa.

I barioni appartenenti all'ottetto erano tutti quelli il cui **spin** è pari a $J = \frac{1}{2}$. Tutti gli altri, infatti, avevano spin $J = \frac{3}{2}$. Anche questi ultimi si possono disporre su una Tavola, usando gli stessi nume-

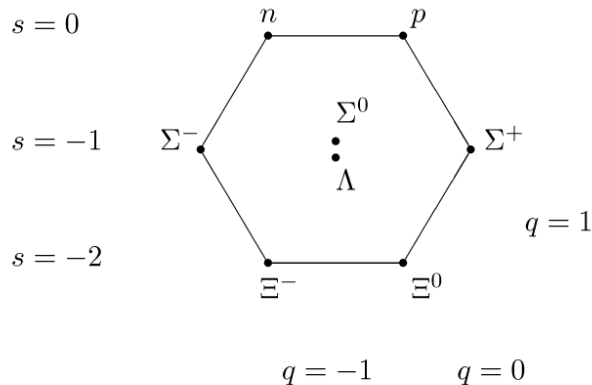


Figura 22.2 L'ottetto di barioni.

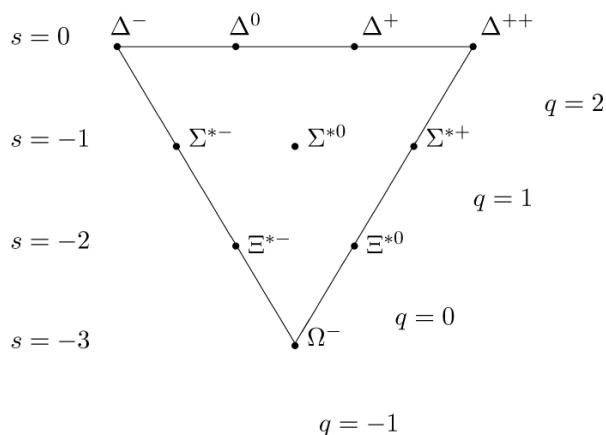


Figura 22.3 Il decupletto di barioni.

ri quantici. L'aspetto della Tavola che se ne ricava è leggermente diverso ed è quello di Fig. 22.3, chiamato il **decupletto di barioni**.

A ben vedere il decupletto di barioni non è molto diverso dall'ottetto. Di fatto, se si eliminano i vertici del triangolo formato dalle particelle che vi si dispongono, la Tavola assume lo stesso aspetto degli ottetti, anche se nella riga superiore compaiono particelle a stranezza nulla e nel mezzo c'è una sola particella, invece che due.

22.2 L'ipotesi dei quark

Ipotizzando che le particelle scoperte non siano elementari, ma a loro volta composte di altre particelle, possiamo pensare che siano il risultato della combinazione di due, tre o più particelle. Con due particelle u e d si possono costruire tre combinazioni diverse: uu , ud e dd (du è evidentemente equivalente a ud). Con tre particelle u , d e s , si possono invece fare 10 combinazioni: uuu , uud , uus , udd , uds , uss , ddd , dds , ssd , sss . Esattamente in numero tale da riprodurre il decupletto di barioni. Ipotizziamo dunque che le particelle del decupletto di barioni siano formate dalla combinazione di tre nuove particelle dette **quark** denominati **up**, **down** e **strange**. Questi **barioni** hanno spin $J = \frac{3}{2}$, quindi potremmo pensare che ciascuno dei tre quark abbia spin $J = \frac{1}{2}$ che, sommandosi, dia luogo a una particella di spin $\frac{3}{2}$. Poiché i barioni hanno stranezza variabile tra 0 e -3 , possiamo attribuire a uno dei tre quark (lo *strange*) stranezza $S = -1$. In questo modo, combinando tre quark s si può ottenere una particella con stranezza $S = -3$. Questa particella (la Ω) deve avere carica elettrica pari a -1 in unità di carica del protone. Perciò i tre quark s devono avere ciascuno carica elettrica pari a $-\frac{1}{3}$ nelle stesse unità. Se è così le due Ξ con stranezza $S = -2$ si devono ottenere dalle combinazioni (che effettivamente sono due) con due quark s : ssu e ssd . Le due Ξ^* hanno carica elettrica -1 e 0 , rispettivamente. Uno dei quark u o d , quindi, deve avere la stessa carica elettrica di s , in modo da produrre insieme a questi una particella con carica elettrica pari a -1 . Il quark d potrebbe dunque essere una particella di spin $\frac{1}{2}$ e carica elettrica $Q = -\frac{1}{3}$ (sempre in unità di carica del protone). La combinazione ssu deve quindi avere carica nulla e questo si può ottenere se si attribuisce al quark u la carica $Q = +\frac{2}{3}$.

Con un quark *strange* si possono fare tre combinazioni che danno luogo a tre particelle con stranezza $S = -1$. Le possibili combinazioni sono uus , uds e dds . Se il quark *up* ha carica elettrica $Q = +\frac{2}{3}$ e il *down* $Q = -\frac{1}{3}$ si vede subito che le combinazioni hanno carica elettrica, rispettivamente, $Q = +\frac{2}{3} + \frac{2}{3} - \frac{1}{3} = 1$, $Q = +\frac{2}{3} - \frac{1}{3} - \frac{1}{3} = 0$ e

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 22.4 La costruzione del decupletto di barioni con il modello a quark [<http://www.youtube.com/watch?v=bGb01UXSXVU>].

$$Q = -\frac{1}{3} - \frac{1}{3} - \frac{1}{3} = -1, \text{ esattamente come le } \Sigma^*.$$

Calcolo Combinatorio

La statistica ci dice che se disponiamo di n elementi che possiamo combinare formando *sequenze* di k di questi elementi, nel caso in cui l'ordine non abbia importanza e tra i k elementi ce ne possono essere di ripetuti, il numero di quelle che si chiamano le **combinazioni con ripetizioni** di n oggetti di **classe** k è

$$C_{nk} = \binom{n+k-1}{k} \quad (22.1)$$

Sostituendo $n = 3$ (il numero di possibili diversi quark che possiamo usare) e $k = 3$ (il numero di quark da usare per ogni sequenza) si ottiene

$$C_{nk} = \binom{5}{3} = \frac{5!}{3!(5-3)!} = \frac{5 \times 4 \times 3 \times 2}{3 \times 2 \times (2)} = 10. \quad (22.2)$$

Usando solo i quark u ed s si ottengono combinazioni prive di stranezza. Ce ne sono quattro: uuu , uud , udd e ddd , che hanno, rispettivamente, carica elettrica $Q = 2$, $Q = 1$, $Q = 0$ e $Q = -1$, proprio come le Δ .

Naturalmente questo non è sufficiente per concludere che effettivamente le cose stiano così, ma è per lo meno un forte indizio. Gell-Mann [30] e Okubo [31] proposero indipendentemente una formula, ricavata da complessi argomenti teorici, che prevedeva la massa delle particelle composte di quark. La

Ω^- , in effetti, fu scoperta dopo (nel 1964) l'ipotesi dei quark e rappresentò un primo grande successo di questa teoria.

Successivamente si moltiplicarono le prove in favore dell'esistenza dei quark, attraverso svariate e complesse misure di sezione d'urto.

22.3 L'ottetto di mesoni

L'ottetto di mesoni, formato di particelle prive di spin, si può realizzare assumendo che i mesoni siano formati da un quark e un antiquark, di spin opposto. Con un quark e un antiquark si possono fare le seguenti combinazioni:

1. $u\bar{u}$, che è una particella di carica elettrica nulla (l'antiquark ha carica elettrica opposta a quella del quark corrispondente) e senza stranezza, identificabile col π^0 o con la η ;
2. anche la combinazione $d\bar{d}$ consente di costruire una particella del tutto simile;
3. $u\bar{d}$, che ha carica elettrica $Q = 1$ e non possiede stranezza, che potrebbe essere il π^+ ;
4. il π^- si potrebbe realizzare con la combinazione $d\bar{u}$, che ha carica $Q = -1$;
5. la prima combinazione con stranezza potrebbe essere $s\bar{u}$, che dà luogo a una particella con carica elettrica $Q = -1$ (il K^-);
6. $s\bar{d}$ produce una particella con carica elettrica nulla e stranezza -1 (il $\bar{K}^0$);
7. $\bar{s}u$ permette di costruire una particella di carica elettrica $Q = +1$ che è identificabile col K^+ ;
8. $\bar{s}d$, invece, realizza una combinazione con carica $Q = 0$ e stranezza $S = +1$: il K^0 ;
9. la combinazione $s\bar{s}$ dà origine a un'altra particella di carica elettrica nulla e priva di stranezza come nei primi due casi.

Come si vede esistono nove combinazioni di coppie quark-antiquark e non otto, come si osserva sperimentalmente. Nulla però impedisce che una particella possa essere composta da più di una combinazione: per esempio le due combinazioni $u\bar{u}$ e $d\bar{d}$ potrebbero dare entrambe luogo a un π^0 se sono indistinguibili. D'altra parte la η , oltre a decadere in due fotoni, come il π^0 , decade anche in tre pioni neutri, quasi il 40 % delle volte. La η , quindi, deve avere un contenuto di quark diverso da quello del π^0 . In generale la massa delle particelle aumenta con l'aumentare della stranezza, quindi il quark s deve essere significativamente più pesante degli altri due. Per questo si può assumere che la η sia il risultato della combinazione $s\bar{s}$, mentre il π^0 possa essere composto da $u\bar{u}$ oppure da $d\bar{d}$ ¹. In questo modo si producono otto diverse combinazioni, esattamente quante se ne osservano sperimentalmente.

22.4 L'ottetto di barioni

Resta da spiegare l'ottetto di [barioni](#). Trattandosi di barioni a spin $J = \frac{1}{2}$, non potendo essere costituiti di un solo quark (che ha carica frazionaria), dobbiamo ritenere che si tratti di combinazioni di tre quark, come nel caso del decupletto. Con due quark, infatti, si possono realizzare combinazioni di spin $J = 0$ (quando gli spin dei due quark sono antiparalleli: $\uparrow\downarrow$) o spin $J = 1$ (quando sono paralleli: $\uparrow\uparrow$). Con tre quark, invece, se due di essi si dispongono in modo da produrre una combinazione con spin $J = 0$, il terzo determina lo spin della particella ($\uparrow\downarrow\uparrow$).

In sostanza l'ottetto di barioni dovrebbe essere formato dalle stesse combinazioni del decupletto, solo che in questo caso gli spin dei tre quark sono, rispettivamente $J = +\frac{1}{2}$, $J = -\frac{1}{2}$ e $J = +\frac{1}{2}$. Ma allora perché sono otto e non dieci? In effetti abbiamo già osservato che all'ottetto, di fatto, mancano le particelle disposte ai vertici del triangolo del decupletto, che sono costituite dalle combinazioni uuu ,

¹La [meccanica quantistica](#) prevede la possibilità che una particella sia una *sovrapposizione* di stati diversi, pertanto il fatto che i pioni siano composti di combinazioni diverse di quark è perfettamente giustificata da questo fatto.

Il Principio di Pauli

Ai piú il Principio di esclusione di Pauli appare come uno strano fenomeno, quasi paranormale, che vieta chissà in quale modo il verificarsi di certe configurazioni. Dopo tutto, cosa impedisce a due elettroni di disporsi nello stato? E come fa un elettrone a *sapere* che si trova nello stato di un altro elettrone?

Il Principio è una diretta conseguenza del fatto che le particelle, in Meccanica Quantistica, si possono descrivere attraverso *funzioni d'onda*: equazioni che hanno la stessa forma matematica di un'onda meccanica. Come è noto due onde possono *interferire* tra loro in modo da risultare in un'onda che può avere un'ampiezza variabile tra un minimo di zero a un massimo pari alla somma delle ampiezze, secondo la *fase* relativa. L'equazione che descrive due particelle identiche di spin semintero è tale da produrre un fenomeno d'interferenza per cui le due onde si annullano a vicenda. È questo fenomeno il responsabile del Principio di esclusione.

In effetti, non ci si dovrebbe stupire tanto del fatto che viga un tale Principio in fisica. In fondo, nessuno si stupisce del fatto che due oggetti non possono stare esattamente nello stesso punto $\vec{x}$ dello spazio, avendo la stessa velocità $\vec{v}$. Il fatto è che in meccanica classica lo stato è definito da posizione e velocità di una particella, mentre in meccanica quantistica queste due variabili non hanno molto senso e lo stato è descritto da energia e momento angolare di una particella. Così il *Principio di esclusione di Pauli* altro non è se non la *traduzione* del principio di impenetrabilità dei corpi che vige in meccanica classica.

ddd e *sss*. Per qualche ragione queste combinazioni devono essere vietate nel caso dell'ottetto.

A pensarci bene la cosa strana non è che manchino queste combinazioni nell'ottetto, ma che siano presenti nel decupletto. Infatti, in *Meccanica Quantistica* vige il *Principio di esclusione di Pauli*, secondo il quale due o piú particelle di spin semintero non

possono mai trovarsi nello stesso stato.

Nelle combinazioni *uuu*, *ddd* e *sss* abbiamo tre particelle dello stesso tipo, nella stessa posizione, con la stessa energia e lo stesso stato di spin. Queste combinazioni sono vietate dal Principio di esclusione e dunque le particelle corrispondenti del decupletto non dovrebbero potersi formare. Per il resto è facile spiegare tutte le altre combinazioni: quelle dell'ottetto di barioni sono le stesse del decupletto, con l'unica differenza di avere i quark in stati di spin tali per cui due di essi hanno spin opposto.

22.5 Quark colorati

L'unica spiegazione che permetteva di giustificare l'esistenza delle combinazioni *uuu*, *ddd* e *sss* consisteva nell'assumere che ogni quark avesse un ulteriore numero quantico, una carica, che permetteva di distinguerlo dall'altro.

A questa *carica* venne attribuito il nome di **colore**. Si può pensare che ogni quark possa esistere in tre diversi stati di carica di colore: **rosso** *R*, **verde** *G* e **blu** *B*. Se la combinazione *uuu* è formata da tre quark di colore diverso, il Principio di esclusione di Pauli non è piú violato e lo stato può esistere. Nel caso dell'ottetto di *barioni* possiamo pensare alla combinazione *uuu* come formata da tre quark di colore diverso *RGB*, tali per cui il quark rosso ha spin $J = +\frac{1}{2}$, quello verde spin $J = -\frac{1}{2}$ e quello blu spin $J = +\frac{1}{2}$. Una tale combinazione si può rappresentare come

$$|uuu \uparrow\rangle = |R, \uparrow\rangle |G, \downarrow\rangle |B, \uparrow\rangle . \quad (22.3)$$

Combinazioni altrettanto valide sono

$$|uuu \uparrow\rangle = |R, \uparrow\rangle |B, \downarrow\rangle |G, \uparrow\rangle \quad (22.4)$$

e

$$|uuu \uparrow\rangle = |G, \uparrow\rangle |R, \downarrow\rangle |B, \uparrow\rangle . \quad (22.5)$$

Poiché tutte e tre queste combinazioni sono possibili dobbiamo considerare la *uuu* di spin $J = \frac{1}{2}$ come una particella formata da tutte e tre queste combinazioni, ciascuna presente nel 33 % circa dei casi. In

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 22.5 Una particella di spin $\frac{1}{2}$ formata da tre quark è una sovrapposizione dei tre possibili stati di spin in cui i quark si possono trovare [http://www.youtube.com/watch?v=PfNP_22tx6A].

altre parole dobbiamo pensare che questa particella abbia un contenuto di quark che si può scrivere, usando una notazione un po' più compatta, come

$$\left|uuu, +\frac{1}{2}\right\rangle = |R \uparrow G \downarrow B \uparrow\rangle + |R \uparrow B \downarrow G \uparrow\rangle + |G \uparrow R \downarrow B \uparrow\rangle \quad (22.6)$$

In altri termini lo stato $|uuu, +\frac{1}{2}\rangle$ si deve considerare come una *sovrapposizione* di stati e non come uno stato formato talvolta da una, talvolta dall'altra combinazione, come prescritto dalla Meccanica Quantistica.

Uno dei modi in cui si può esprimere il **Principio di Pauli** consiste nel dire che, qualora scambiando due particelle si ottenga uno stato equivalente a quello iniziale, la combinazione è vietata. Nel nostro caso, se scambiamo di posto la seconda e la terza combinazione otteniamo comunque uno stato identico al precedente e così nel caso di tutti i possibili scambi che possiamo pensare di fare. Queste combinazioni sono dunque vietate dal Principio di Pauli.

La combinazione con spin $J = \frac{3}{2}$

$$|\Delta^{++}\rangle = |R \uparrow G \uparrow B \uparrow\rangle \quad (22.7)$$

è una sola e non è vietata perché i tre quark hanno colore diverso e sono perciò distinguibili (la combinazione $|R \uparrow B \uparrow G \uparrow\rangle$ è la stessa perché i quark B e G si devono intendere nella stessa posizione, in un modo che non possiamo rappresentare sulla carta).

L'introduzione della carica di colore permette così di spiegare le diverse configurazioni osservate, ma

introduce un elemento che farebbe pensare alla possibilità di costruire molte altre combinazioni, come, ad esempio, uds con i quark u e d dello stesso colore e il quark s di colore diverso. Oltre alla combinazione di colori RGB , dunque, potrebbe esserci una combinazione con gli stessi numeri quantici della Σ^0 , ma fatta con i colori RRG o GGR o BBG , etc.. Di combinazioni così ce ne sono diverse e dovrebbero ciascuna dar luogo a particelle che dovrebbero comunque manifestare in qualche modo la carica di colore in eccesso.

In effetti possiamo pensare che la combinazione dei tre colori generi una carica di colore neutra (si dice che la particella è **bianca** o **incolora**), ma nel caso della combinazione RRG si dovrebbe poter osservare questa carica di colore come non neutra e questo dovrebbe produrre effetti sulle interazioni osservate. Se infatti la carica di colore esiste, deve essere associata a una qualche interazione per la quale l'intensità dipende da questa carica (come la carica elettrica, associata alle interazioni elettromagnetiche, ne determina l'intensità), perciò una particella formata da due quark rossi e uno verde dovrebbe presentare un *eccesso di carica rossa* che dovrebbe potersi manifestare attraverso interazioni più o meno intense rispetto a quelle subite da una particella *bianca*.

Non avendo mai osservato alcuna differenza nelle interazioni delle Σ^0 (né delle altre particelle) si pensò che dovessero essere possibili solo combinazioni di colori tali da formare una particella **priva di colore** o bianca. In questo modo le uniche combinazioni possibili sono quelle osservate. Anche nel caso dei mesoni, la combinazione di un quark con un antiquark è bianca, perché se un quark porta una *carica rossa*, l'antiquark ne porta una **antirossa**.

La regola secondo la quale tutte le particelle devono essere bianche permette anche di spiegare come mai, nonostante una lunga serie di tentativi, non furono mai osservati quark liberi. Si potrebbe infatti pensare che, colpendo un protone (formato da due quark *up* e da un quark *down*) con una particella con energia sufficiente, se ne dovrebbero poter estrarre i quark costituenti perché nell'urto i legami che tengono insieme i quark potrebbero spezzarsi. In effetti

si potrebbe pensare che i quark sono tenuti insieme nel protone (e nelle altre particelle) dalle **interazioni forti** e che quelle che si osservano nei nuclei, che tengono insieme neutroni e protoni, siano il residuo di queste forze visto attraverso la schermatura dei costituenti: una specie di **forza di Van der Waals** forte. Disponendo di sufficiente energia si dovrebbe poter vincere l'attrazione prodotta da queste forze e liberare così i quark che si dovrebbero poter osservare in esperimenti di questo genere come tracce elettricamente cariche, ma di carica frazionaria.

Il fatto che le particelle fisiche (quelle osservabili sperimentalmente) debbano essere bianche permette di spiegare l'assenza di questi eventi. È impossibile per i quark essere prodotti liberi. È quello che si chiama il fenomeno del **confinamento**.

Naturalmente esistono numerose prove dell'esistenza dei quark, sebbene non siano mai stati osservati liberi. L'esistenza di queste tre particelle e della carica di colore si prova attraverso misure che non è il caso di illustrare in questa sede, ma che dimostrano inequivocabilmente che le particelle fisiche sono formate da combinazioni di quark tenuti insieme dalla forza forte, che impedisce loro di *uscire* dalle particelle costituenti.

Non si vedono, ma si contano!

La Fisica è una scienza sperimentale che richiede di poter **osservare** i fenomeni di cui tratta. Sembrerebbe dunque di poter affermare che, giacché i quark sono intrinsecamente inosservabili dal momento che non si possono produrre liberi, non dovrebbero poter essere soggetti d'indagine per un fisico. In realtà, quando un fisico pronuncia la parola **osservare** intende **misurare** (non è infrequente che per i fisici il significato delle parole sia diverso da quello loro comunemente attribuito: è un fenomeno che si deve alla possibilità per la lingua comunemente parlata di essere interpretata in maniera ben più ambigua di quanto non sia necessario in fisica o in matematica).

All'inizio del secolo scorso non erano pochi i fisici che ancora non credevano all'esistenza degli **atomi**. L'argomento per confutarne l'esistenza era molto efficace, in effetti: gli atomi non si possono vedere, quindi non esistono! Fu **Jean Baptiste Perrin** a confutare questa tesi, eseguendo molte misure del **Numero di Avogadro** che davano tutto lo stesso risultato, basandosi sui lavori del solito **Einstein** sul **moto browniano**. Perrin, che per quei lavori vinse il Premio Nobel, osservò che è vero che gli atomi non si possono vedere perché sono più piccoli della lunghezza d'onda della luce, ma si possono **contare**, quindi se ne può misurare il loro numero, pertanto esistono.

Allo stesso modo, è vero che i quark non si possono osservare, ma se ne possono misurare gli effetti della loro esistenza in molti modi (non solo quelli illustrati qui). Dunque, benché sia intrinsecamente impossibile produrli singolarmente e **osservarli** in qualche maniera, possiamo sicuramente dire che esistono.

Il Modello Standard

Tipo	Carica elettrica	Famiglia		
	Q	I	II	III
Quark	$+\frac{2}{3}$	u	c	t
	$-\frac{1}{3}$	d	s	b
Leptoni	-1	e	μ	τ
	0	ν_e	ν_μ	ν_τ

Tavola 23.1 Il Modello Standard delle particelle elementari prevede l'esistenza di sei quark e sei leptoni, oltre alle rispettive antiparticelle, come facenti parte dei costituenti della materia dell'Universo.

Il **Modello a Quark** aveva ricondotto la teoria della fisica delle particelle a una condizione di semplicità. Ora bastavano solo tre quark per spiegare la materia, insieme agli elettroni. Restava da capire il ruolo del muone e dei due neutrini, ma era stato fatto un grosso passo in avanti.

23.1 I costituenti della materia

Col passare del tempo si scoprirono nuove particelle e il quadro si complicò di nuovo anche se, diventando più *simmetrico* acquistò maggiore solidità. Interpretando tutti i dati sperimentali fino ad ora conosciuti, sappiamo che esistono sei quark: oltre a u , d e s esistono il **charm** c , di carica $+2/3$, il **top** t , anch'esso di carica $+2/3$ e il **bottom** o **beauty** b con carica uguale a quella del down e dello strange.

Esistono anche sei leptoni: oltre all'elettrone, al muone e ai rispettivi neutrini, esiste infatti il **tau** con il rispettivo neutrino ν_τ . Esiste dunque una per-

fetta simmetria tra leptoni e quark, che possiamo dividere in **famiglie**. La prima famiglia, che contiene le particelle più leggere, è formata dai quark *up* e *down*, dall'elettrone e dal suo neutrino (con lo stesso numero leptonico). La seconda famiglia, oltre al *charm* e allo *strange*, include il muone e il neutrino muonico, mentre della terza fanno parte il *top* e il *bottom*, il τ e il neutrino ν_τ .

La materia ordinaria è formata solo di particelle della prima famiglia. Le altre famiglie sono di fatto copie più pesanti della prima. A oggi non sappiamo ancora perché esistano tre famiglie di particelle. In effetti basterebbe la prima per spiegare la composizione dell'intero Universo.

Tutti i fenomeni osservati si spiegano alla luce di questo Modello. Il decadimento del neutrone, ad esempio, s'interpreta come il decadimento di un quark d che decade secondo la reazione

$$d \rightarrow u + e^- + \bar{\nu}_e \tag{23.1}$$

così che da una particella formata da due quark *down* e un quark *up* si formi una particella composta di due quark *up* e un quark *down* (il protone) con l'emissione di una coppia di leptoni con **numero leptonico** opposto.

La conservazione del numero barionico si spiega dunque con l'impossibilità di sopprimere un quark. Un quark può trasformarsi in un altro oppure generare una coppia quark-antiquark. Non esistono regole di conservazione per i mesoni perché sono composti di quark e antiquark: se uno di questi si trasforma il numero di quark non cambia. La conservazione del numero leptonico, invece, indica che le trasformazioni dei leptoni possono avvenire solo all'interno dello stesso **doppietto**, che non si posso-

no né sopprimere né creare leptoni, se non in coppie leptone–antileptone.

Quark e leptoni possono avere una carica elettrica e, se ce l'hanno, sono soggetti a interazione elettromagnetica. Sia quark che leptoni sono anche soggetti all'[interazione debole](#). Al contrario dei quark, i leptoni non subiscono l'interazione forte. Per i leptoni è come se quest'interazione non esistesse.

In definitiva nell'Universo possono esistere sei tipi di quark e altrettanti leptoni. Oltre, naturalmente, alle rispettive antiparticelle. È con questi ingredienti che si costruiscono tutte le particelle osservate sperimentalmente. In passato le particelle più pesanti di quelle della prima famiglia dovevano essere più abbondanti di quelle presenti oggi, perché l'Universo, più compatto, era molto più caldo e le particelle in esso contenute avevano energie molto più elevate di quelle odierne. Le collisioni tra le particelle presenti potevano dunque generare Λ , Σ , Ξ , K , etc. Col tempo l'Universo si è raffreddato e l'energia delle particelle non è più stata sufficiente a produrne di nuove. Quelle esistenti hanno cominciato a decadere e oggi sono rimasti praticamente solo protoni e neutroni. Le altre particelle possiamo osservarle solo nelle rare collisioni ad alta energia dei raggi cosmici oppure in laboratorio. In un certo senso, dunque, lo studio della fisica delle particelle con gli acceleratori è anche lo studio dell'evoluzione dell'Universo e gli stessi acceleratori sono come **macchine del tempo** per tornare a epoche remotissime, quando l'Universo era molto giovane e aveva un'età compresa tra qualche frazione di secondo fino a qualche minuto.

Fisicast

In [Fisicast](#) abbiamo parlato di particelle elementari nei podcast:

Il microscopico zoo delle particelle elementari

<http://www.radioscienza.it/2013/09/22/il-microscopico-zoo-delle-particelle-elementari/>

Vedere le particelle elementari [http://](http://www.radioscienza.it/2014/01/20/vedere-le-particelle-elementari/)

www.radioscienza.it/2014/01/20/vedere-le-particelle-elementari/

Campi e Particelle

PREREQUISITI: dinamica elementare, e energia meccanica, concetto di potenziale di un campo

La Fisica delle Particelle non consiste soltanto nell'individuare i costituenti elementari della materia: è soprattutto lo studio delle interazioni cui tali costituenti sono soggetti. Di fatto, dunque, la Fisica delle Particelle è la disciplina che studia le **forze fondamentali** che si manifestano nell'Universo e che determinano, in fin dei conti, il comportamento degli oggetti macroscopici (dagli atomi, ai batteri, agli esseri viventi evoluti come noi, ai pianeti e alle galassie).

Le forze fondamentali sono quelle che non si possono ricondurre all'azione combinata di altre forze. Le forze elastiche, ad esempio, sono il risultato delle interazioni tra i costituenti dei materiali di cui sono composte le molle, che sono di natura elettromagnetica. Al contrario, esistono alcune forze che sembrano appartenere a classi diverse, non riconducibili l'una all'altra e che per questo dobbiamo considerare come fondamentali (almeno fino a quando non scopriremo il contrario).

24.1 Le forze fondamentali

A oggi conosciamo quattro forze che possiamo considerare fondamentali: la **gravità** è forse la più conosciuta. È quella responsabile della caduta degli oggetti sulla Terra e del fatto che i pianeti orbitano attorno al Sole (e le stelle orbitano attorno al centro della Galassia). Anche le forze **elettromagnetiche** sono abbastanza note: sono quelle che danno origine ai fenomeni di tipo elettrico o magnetico. Le usiamo

pesantemente per **produrre l'elettricità** che ci serve ad alimentare tutti i nostri apparati elettronici.

Ci sono altre due forze meno note: la **forza debole** e quella **forte**. Quest'ultima è responsabile del fatto che i nuclei atomici stanno insieme nonostante le forze di natura elettromagnetica producano una forte repulsione tra i protoni. In assenza di questa forza due protoni in un nucleo (a distanze quindi dell'ordine di $1 - 2$ fm, cioè di $1 - 2 \times 10^{-15}$ m) si respingerebbero con una forza immensa pari a

$$F = \frac{1}{4\pi\epsilon_0} \frac{q^2}{r^2} \simeq 9 \times 10^9 \frac{4 \times 10^{-38}}{10^{-30}} = 360 \text{ N} \quad (24.1)$$

che è la forza peso di un ragazzino. Affinché i protoni non si allontanino l'uno dall'altro deve esistere una forza molto più intensa che deve trattenerli all'interno del nucleo: la forza forte, appunto. Questa forza deve avere un raggio d'azione piuttosto piccolo. In effetti i protoni liberi non sembrano essere soggetti a forze attrattive così forti da parte di altre cariche positive, quindi la forza forte deve *spegnersi* abbastanza rapidamente.

La forza debole invece è quella che permette le interazioni tra i neutrini e le altre particelle e che provoca i decadimenti radioattivi dei nuclei atomici. È importante comprendere che le forze non causano solo il movimento (come si tende a credere per via della nota **equazione di Newton** secondo la quale $F = ma$), ma in generale provocano il cambiamento dello stato di un oggetto. Lo stato dei *punti materiali* è perfettamente determinato quando se ne conoscono posizione e velocità: se cambia una di queste due grandezze dev'essere intervenuta una forza. In sistema complesso lo stato è determinato

dalla sua temperatura, dal volume e dalla pressione. Un atomo non soggetto a forze non può cambiare il suo stato se non per mezzo dell'applicazione di una forza. In seguito all'intervento della forza debole il nucleo di quell'atomo può trasformarsi in un nucleo di specie diversa.

24.2 Una rivisitazione del concetto di energia

In Fisica Classica trattiamo spesso particelle come punti materiali soggette a forze di varia natura. Lo stato delle particelle è determinato quando se ne conoscano posizione e velocità. Se i corpi considerati non hanno massa costante (è il caso di un'automobile che man mano che procede consuma carburante), anche la massa determina in qualche modo lo stato del sistema. In certi casi può essere importante la temperatura, etc..

Se un sistema composto da una o più particelle è **isolato** (non può cioè interagire con nulla) possiamo sempre definire una grandezza fisica E , funzione dello stato delle particelle che compongono il sistema, che ha la proprietà secondo cui

$$\Delta E = 0, \quad (24.2)$$

cioè che la sua variazione è nulla. È del tutto evidente che, data una qualunque legge fisica, come $F = ma$, si possa definire $E = F - ma$ e quindi $E = 0$, pertanto anche $\Delta E = E(t + \delta) - E(t) = 0$, avendo indicato con $E(t)$ la grandezza fisica E misurata al tempo t . In questo non c'è nulla d'interessante. Se però riusciamo a individuare combinazioni utili di grandezze fisiche per cui vale quanto sopra possiamo imparare qualcosa di nuovo. Consideriamo, ad esempio, l'**energia meccanica** di una particella di massa m posta a una quota h dal suolo, che si muove con velocità v :

$$E = \frac{1}{2}mv^2 + mgh. \quad (24.3)$$

Questa quantità dipende solamente dallo stato che la particella assume in ogni istante di tempo ed è

costante: $E(t) = E_0 = \text{const.}$ Se è costante, la sua variazione nel tempo è nulla. Indicando col pedice f le grandezze nello stato finale e col pedice i quelle nello stato iniziale possiamo scrivere che

$$\frac{\Delta E}{\Delta t} = 0 = \frac{1}{2\Delta t}mv_f^2 + mg\frac{h_f}{\Delta t} - \frac{1}{2\Delta t}mv_i^2 - mg\frac{h_i}{\Delta t}. \quad (24.4)$$

Raccogliendo i termini simili si ottiene

$$0 = \frac{1}{2}m\frac{(v_f^2 - v_i^2)}{\Delta t} + mg\frac{(h_f - h_i)}{\Delta t} \quad (24.5)$$

e osservando che $v_f^2 - v_i^2 = (v_f - v_i)(v_f + v_i)$ possiamo scrivere

$$\frac{1}{2}m(v_f + v_i)\frac{(v_f - v_i)}{\Delta t} = -mg\frac{(h_f - h_i)}{\Delta t}. \quad (24.6)$$

Ora, $(h_f - h_i)/\Delta t$ non è altro che lo spostamento subito dalla particella nell'intervallo Δt diviso per questo stesso intervallo, quindi non è altro che la velocità media della particella v . Ma anche $(v_f + v_i)/2 = v$; inoltre $(v_f - v_i)/\Delta t = a$ non è altro che la variazione della velocità nell'unità di tempo, cioè l'accelerazione della particella, perciò, sostituendo si trova che

$$mva = -mgv \quad (24.7)$$

da cui, dividendo tutto per v si ottiene che la massa della particella per la sua accelerazione a è pari al suo peso mg , come previsto dalla legge di Newton. In generale non è difficile rendersi conto che tutta la dinamica dei corpi è contenuta in un principio fondamentale che è quello della **conservazione dell'energia**.

È il fatto che l'energia si conserva a far sì che $F = ma$. Quest'ultima, in altri termini, è una conseguenza del principio di conservazione dell'energia. Il che non significa che nell'Universo tutto deve restare immobile. Una pallina sollevata a un'altezza h possiede una certa quantità d'energia. Posso aumentare l'**energia cinetica** $\frac{1}{2}mv^2$ della pallina, ma poiché l'energia totale si deve conservare, questo può solo avvenire in seguito a una diminuzione dell'**energia**

potenziale mgh . Il che implica che sia cambiato lo stato della particella: l'altezza h non è più la stessa. Noi imputiamo questo cambiamento di stato alla presenza di una **forza** che altro non è se non una misura dell'entità del cambiamento.

Osserviamo che la derivata dell'energia potenziale U rispetto alla posizione h è $dU/dh = mg$, che è proprio la forza con la quale il corpo cade. Possiamo dunque dire che la dinamica dei sistemi è determinata unicamente da una funzione scalare U , che è funzione solo delle coordinate, definita in tutto lo spazio. La variazione del valore di questa funzione in un punto dello spazio determina la comparsa di una forza e quindi la variazione della velocità della particella che si trova in quel punto.

Quando scriviamo $U = mgh$ per indicare l'energia potenziale gravitazionale sappiamo bene che questa è definita a meno di una costante. Potremmo benissimo definire $U = mgh + C$ con C costante e non cambierebbe nulla: quello che conta per determinare la dinamica del sistema infatti non è l'energia, ma la sua derivata (la sua variazione). Un altro modo di vedere la stessa cosa è dire che possiamo scegliere come vogliamo il livello al quale $U = 0$. Ancora una volta questo dipende dal fatto che non è possibile misurare l'energia, ma solo **differenze di energia**.

L'energia potenziale gravitazionale si definisce come il **lavoro** fatto dalle forze gravitazionali cambiato di segno:

$$\Delta U = U(h) - U(0) = - \int_h^0 mg \, dh = mgh. \quad (24.8)$$

Benché di norma non si faccia, si potrebbe definire il **potenziale gravitazionale** V come l'energia potenziale per unità di massa:

$$V = V(h) = \frac{U}{m} = gh \quad (24.9)$$

e l'energia potenziale gravitazionale per un corpo di massa m alla quota h si potrebbe esprimere come $mV(h)$. Quando si studiano le **interazioni gravitazionali** s'impara però che il potenziale gravitazionale è un'altra cosa: è il lavoro per portare una massa unitaria dall'infinito alla posizione nella quale si

deve valutare il potenziale, che dista r dalla sorgente del campo gravitazionale. Attorno alla Terra, dunque, il potenziale gravitazionale sarebbe

$$V = V(r) = \int_{\infty}^r G \frac{M}{r'^2} dr' = -G \frac{M}{r} \Big|_{\infty}^r = -G \frac{M}{r}. \quad (24.10)$$

Le due espressioni del potenziale hanno in comune il fatto di essere entrambe definite come il lavoro fatto dalle forze gravitazionali per portare una massa unitaria da un punto all'altro, cambiato di segno, e di crescere al crescere della distanza dalla sorgente del campo, ma appaiono decisamente diverse l'una dall'altra. Se la forza che tiene insieme il sistema solare è la stessa che fa cadere i corpi sulla Terra, come si spiega che il potenziale dell'una è diverso dal potenziale dell'altra? In realtà non c'è alcuna differenza: quando scriviamo che il potenziale gravitazionale è $V = gh$ stiamo semplicemente assumendo che g sia costante, ma questa è solo un'approssimazione. In effetti sappiamo che il potenziale della Terra, in prossimità della sua superficie, varia sí, ma di pochissimo, perché $h \ll r_0$, dove r_0 è il raggio della Terra. Sostituendo a g la sua espressione

$$g = G \frac{M}{r_0^2}, \quad (24.11)$$

scrivere $V = gh$ equivale a scrivere

$$V = G \frac{M}{r_0^2} h. \quad (24.12)$$

Scriviamo ora l'espressione esatta di V , come data nell'equazione (24.10), approssimandola con una retta (lo possiamo fare se $r_0 + h \simeq r_0$, cioè quando $h \ll r_0$):

$$V(r) = -G \frac{M}{r} = -G \frac{M}{R_0 + h} \simeq -G \frac{M}{r_0} \left(1 - \frac{h}{r_0} \right). \quad (24.13)$$

Il primo addendo di questa somma è una costante irrilevante (contano solo le differenze di potenziale). Il secondo vale

$$G \frac{M}{r_0^2} h = gh. \quad (24.14)$$

Come si vede, scrivere $U = mgh$ equivale a considerare le distanze h come molto piccole rispetto al raggio terrestre. La vera espressione di U dovrebbe essere, in effetti $U = mV$ dove V è dato dall'equazione (24.10). Si tratta di un risultato completamente generale. Ogni funzione si può approssimare, nelle vicinanze di un punto x_0 , con un polinomio. Maggiore è il grado del polinomio, migliore è l'approssimazione con la quale si approssima il valore della funzione. In generale quindi

$$V(r) = -G \frac{M}{r} = -G \frac{M}{R_0 + h} \simeq -G \frac{M}{r_0} \left(1 - \frac{h}{r_0} + \frac{h^2}{2r_0^2} \right) \quad (24.15)$$

Non è difficile convincersi che lo stesso accade per tutte le altre forze fondamentali. Se calcoliamo il lavoro fatto dalle forze elettriche cambiato di segno otteniamo l'espressione dell'energia potenziale elettrostatica, che è pari al potenziale elettrostatico moltiplicato per la carica elettrica.

Ne concludiamo che, in effetti, potremmo riassumere tutta la fisica in un'unica legge: **l'energia totale dell'Universo è e deve rimanere costante**. Se cambia l'energia in una regione dell'Universo, deve avvenire qualche cambiamento in un'altra regione tale per cui la somma algebrica delle variazioni di energia sia nulla. Se non ci fossero interazioni l'energia dell'Universo sarebbe data dalla somma delle energie cinetiche di tutte le particelle in esso contenute, che non cambierebbe mai perché in assenza di interazioni non ci sarebbero accelerazioni e dunque le velocità delle particelle non potrebbero cambiare. In presenza di interazioni possiamo associare a ogni punto dell'Universo una funzione scalare delle coordinate che chiamiamo energia potenziale la cui variazione deve essere compensata da una variazione opposta dell'energia cinetica. Il *gradiente*, cioè la rapidità con la quale cambia, dell'energia potenziale rappresenta la forza che si osserva sperimentalmente in virtù di questo principio.

24.3 L'energia delle interazioni tra particelle

L'energia potenziale dunque è l'energia determinata dalla presenza di interazioni: in presenza di due o più corpi interagenti si può definire questa quantità. I corpi devono evidentemente essere almeno due altrimenti non avremmo nessun tipo d'interazione.

Mettiamoci nel caso più semplice possibile di due sole particelle elementari che interagiscono. Una delle due particelle la possiamo considerare sorgente di un campo di forze, subito dall'altra che indichiamo come *particella di prova* o viceversa (ricordate sempre il terzo principio della dinamica). La posizione dell'una rispetto all'altra è dunque perfettamente determinata dalla distanza $\mathbf{r}$ tra le due particelle. L'energia potenziale che possiamo associare alla particella di prova deve dunque essere una funzione scalare di $\mathbf{r}$: $U = U(\mathbf{r})$. Evidentemente deve dipendere dal tipo di particelle interagenti: una cosa è se le particelle hanno carica elettrica, una cosa è se non ce l'hanno. Se indichiamo con la lettera greca ϕ l'insieme delle caratteristiche delle particelle che determinano l'apparire dell'interazione possiamo scrivere che U deve essere funzione anche di questo insieme per ciascuna delle due particelle: $u = U(\mathbf{r}, \phi_1, \phi_2)$. Ovviamente, questa energia dipenderà dal tipo di campo di forze che la particella sorgente genera, che indichiamo con la lettera A : una cosa, ad esempio, è se le particelle cariche sono in quiete e una cosa è se sono in moto. Nel primo caso interagiscono elettrostaticamente, attraverso un campo elettrico; nel secondo caso sarà presente anche un campo magnetico. In definitiva $U = U(\mathbf{r}, \phi_1, \phi_2, A)$. La distanza $\mathbf{r}$ dipende unicamente dalle coordinate di ϕ_1 e ϕ_2 , che rappresentano le due particelle interagenti, perciò l'energia dipende dalle coordinate e dunque

$$U = U(\phi_1, \phi_2, A(\mathbf{x}_1, \mathbf{x}_2)) \quad (24.16)$$

L'espressione di U può essere complicata a piacere, ma purché le dimensioni fisiche delle cose da cui dipende siano quelle opportune, possiamo sempre approssimarla con un polinomio $U = U_0 + U_1 y + U_2 y^2 + \dots$. La costante U_0 è irrilevante, perché quello che

conta è solo la differenza di energia, quindi possiamo sempre porre $U_0 = 0$. La variabile y del polinomio sarà una combinazione delle caratteristiche di ϕ_1 , ϕ_2 e A . Nel caso più semplice avremo

$$y = \phi_1 A(\mathbf{x}_1, \mathbf{x}_1) \phi_2 \quad (24.17)$$

cioè un semplice prodotto delle caratteristiche. In questo modo se uno dei $\phi_i(\mathbf{x}_i) = 0$, vale a dire la particella corrispondente è assente, $U = 0$. Lo stesso accade se $A = 0$: il campo è assente e non ci sarà interazione.

In generale, dunque, U sarà una funzione qualunque $f(y)$ di questo prodotto, ma sarà sempre esprimibile, almeno in un intorno del punto d'interesse, come un polinomio: al limite come una retta. Un esempio chiarirà meglio la questione: consideriamo sempre le interazioni gravitazionali che sono quelle che conosciamo meglio. L'energia di una particella di massa m nel campo di una particella di massa M si scrive come

$$U = G \frac{Mm}{r}. \quad (24.18)$$

U è funzione delle caratteristiche che determinano, per ciascuna delle due particelle, l'interazione ($\phi_1 = M$ e $\phi_2 = m$) e dalla forma del potenziale del campo che possiamo indicare genericamente come $A = G/r$ (il potenziale prodotto da una particella ϕ si ottiene moltiplicando ϕ per A). Se invece di scrivere l'energia per esteso, ne scriviamo un'espressione approssimata, trascurando la presenza della costante irrilevante, scriveremmo

$$U \simeq \phi_1 A \phi_2 = M \frac{G}{r} m \quad (24.19)$$

che è proprio l'espressione esatta dell'energia potenziale (questo significa che, nel caso in esame, $U_1 = 1$ e $U_{i \neq 1} = 0$). Dunque non importa quanto sia complicata la funzione che definisce la *vera* energia potenziale di tutte le particelle dell'Universo: possiamo sempre restringerci a considerare due particelle sufficientemente vicine da interagire in modo da produrre un'energia potenziale pari a questo prodotto di *caratteristiche* $\phi_1 A \phi_2$.

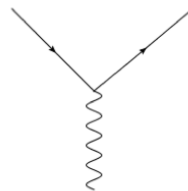


Figura 24.1 Il diagramma di Feynman più semplice, che rappresenta l'interazione tra un campo e due particelle.

Graficamente possiamo rappresentare l'equazione che definisce il primo termine dello **sviluppo polinomiale** dell'energia in questo modo: disegniamo una freccia ogni volta che compare un fattore ϕ_i ; le due frecce che inevitabilmente dovranno comparire le disegniamo contigue, per indicare che l'interazione tra queste due particelle avviene in un preciso punto dello spazio che è quello in cui la prima freccia converge e da cui la seconda emerge. Da questo stesso punto disegniamo poi una linea ondulata che rappresenta il campo che produce l'interazione tra le due particelle, come nella Figura 24.1. Il diagramma di questa figura si chiama **diagramma di Feynman**¹.

I diagrammi di Feynman sono un vero e proprio strumento di calcolo usato dai Fisici Teorici per calcolare le probabilità che avvengano certi fenomeni e sono troppo complessi per poter essere insegnati a questo livello. Possiamo però provare a darne una *libera* interpretazione che ne chiarisce il senso vero.

In questo diagramma abbiamo due particelle e un campo nello stesso preciso punto dello spazio: il **vertice** dell'interazione, costituito dal punto in cui convergono le tre linee. Questo non è ragionevole: due particelle non possono stare nello stesso identico punto! Si potrebbe viceversa interpretare questo diagramma come segue: una particella di materia proviene da sinistra; a un certo punto, quando si trova nel vertice, *produce* un campo, che è rappresentato dalla linea ondulata. In questa visione il campo

¹Dal nome del fisico **Richard Feynman** che li inventò.

non è qualcosa di permanente nello spazio, ma qualcosa che viene *emesso* dalla particella in continuazione. In seguito all'emissione del campo, che è un processo simile a quello in cui si *lancia* qualcosa da un oggetto in corsa, per effetto della conservazione della quantità di moto, la particella cambia leggermente direzione. In altre parole, quello che abbiamo appena descritto è un processo nel quale una singola particella emette un campo che non influenza il moto di nessun'altra particella nelle vicinanze. In seguito a questa emissione la particella cambia direzione mentre il campo continua a viaggiare fino a distanze infinite. Anche questo non appare molto ragionevole: una sola particella magari emette un campo, ma se non ci sono particelle nelle vicinanze non potremo mai saperlo, dal momento che non c'è modo d'interagire con essa e carpirne qualche informazione.

In effetti, nella teoria di Feynman, usando gli opportuni *valori* per ϕ_1 , ϕ_2 e A si trova che il prodotto $U = \phi_1 A \phi_2$ è nullo. Questo indica che la probabilità che accada questo fenomeno è nulla e quindi questo diagramma non può contribuire all'energia dell'Universo. Di conseguenza il termine $U_1 \phi_1 A \phi_2 = 0$. Non è quello che succede all'energia potenziale gravitazionale, ma solo perché in quel caso stiamo considerando corpi costituiti di moltissime particelle!

Il primo termine dello sviluppo che potrebbe essere non nullo è dunque il termine di grado due:

$$U \simeq U_2 (\phi_1 A \phi_2)^2. \quad (24.20)$$

Dal punto di vista grafico possiamo pensare a questo termine come composto dal prodotto di due diagrammi come quelli di Figura 24.1, opportunamente uniti. Un modo per farlo è quello di unire il diagramma della Figura 24.1 con un altro identico attraverso la linea ondulata: il campo. Come nella Figura 24.2.

In questo diagramma due particelle si muovono da sinistra verso destra avvicinandosi. A un certo istante una delle due (quella di sopra, per esempio) emette un campo e cambia direzione. L'altra (quella in basso) *raccoglie*, per così dire, il campo e, come nel caso in cui qualcuno raccolga un oggetto

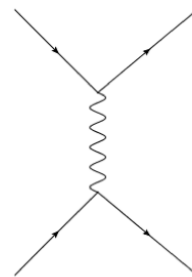


Figura 24.2 L'interazione tra due particelle come descritta da un diagramma di Feynman.

pesante lanciato da qualcun altro, cambia direzione anche lei. Naturalmente potremmo invertire il ruolo del *lanciatore* e del *raccoglitore*: non cambierebbe nulla. Abbiamo appena descritto un processo nel quale due particelle si avvicinano e si respingono l'una con l'altra: non si può dire chi respinge quale: si respingono a vicenda. L'una si può considerare produttrice del campo e l'altra quella che lo subisce o viceversa. L'effetto è lo stesso.

Calcolando questo diagramma con le regole di Feynman si trova in effetti un valore di probabilità diverso da zero. Se poi si confronta questa probabilità d'interazione con quella misurata sperimentalmente, usando le corrette espressioni per ϕ_1 , ϕ_2 ed A , si trova che queste praticamente coincidono! In pratica con le regole di Feynman si può calcolare la probabilità che, ad esempio, un elettrone sia diffuso da un altro elettrone a un certo angolo, avendo l'uno una certa quantità di moto e l'altro, per esempio, sia fermo. Si può quindi misurare questa probabilità misurando la frequenza con la quale elettroni con la quantità di moto scelta sono diffusi da elettroni fermi all'angolo desiderato. Il risultato è un sorprendente accordo tra teoria ed esperimento.

Diremo, allora, che l'interazione tra due particelle funziona così: evidentemente le particelle di materia possono *emettere* dei campi. Se nelle vicinanze si trova una particella in grado di assorbire tale campo, lo fa e in questo modo si manifesta l'interazione. In pratica le due particelle interagiscono perché si

scambiano qualcosa che chiamiamo **campo**, ma che potremmo pensare come a una particella prodotta dalla prima e raccolta dalla seconda. Questo campo è dunque rappresentabile come una **particella mediatrice** di forza che è scambiata tra le particelle di materia che interagiscono. Noi non possiamo *osservare* direttamente questo processo: si tratta soltanto di un'astrazione matematica. Ma questa astrazione rappresenta bene la realtà sperimentale e pertanto siamo autorizzati a pensare che le cose vadano effettivamente così: non importa se non vanno davvero così: la Fisica è una scienza sperimentale e come tale descrive le osservazioni. Le descrive in termini matematici. L'interpretazione che diamo delle equazioni in termini di *oggetti* è del tutto arbitraria, anche se funzionale. Del resto sappiamo bene che i pianeti non sono dei punti, ciò non di meno si possono rappresentare così nella nostra testa quando ne consideriamo il moto descritto dalle equazioni di Newton, che sono l'unica *cosa reale*.

In questo modo il campo elettromagnetico diventa una particella (che chiameremo **fotone**) che è scambiata tra due elettroni e ne provoca la repulsione. Ma come si spiega invece l'attrazione tra un elettrone e un protone? Proviamo a insistere con questa interpretazione prendendo un protone inizialmente fermo. Questo, a un certo punto dovrebbe emettere un fotone e quindi dovrebbe muoversi nel verso opposto a quello nel quale si muove il fotone. Se giunge nelle vicinanze un elettrone che si muove parallelamente al fotone in direzione di quest'ultimo, si scontra con questo e riceve una spinta all'indietro. Apparentemente, dunque, il processo non funziona: il protone e l'elettrone si respingerebbero invece di attrarsi. Ma non dobbiamo dimenticare quanto abbiamo detto sopra! La nostra descrizione qualitativa del processo è una libera interpretazione delle equazioni che sono l'unica descrizione valida della realtà. Basta cambiare il segno dell'energia per provocare un'interazione attrattiva, quindi evidentemente lo stesso processo deve poter descrivere anche questo tipo d'interazioni. Possiamo continuare a immaginare il processo come lo scambio di qualcosa: basta cambiare prospettiva! Pensiamo a due giocatori che si scambiano delle clavette: le clavette viaggiano in

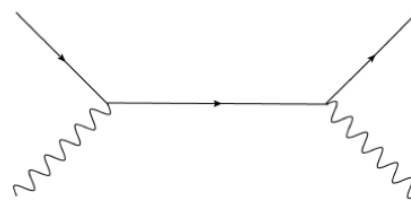


Figura 24.3 L'interazione tra una particella e un campo.

continuazione dall'uno all'altro e viceversa. Fino a quando i due giocatori si scambiano questi oggetti sono *costretti* in qualche modo a restare vicini. Sono dunque *attratti* l'uno dall'altro. Se smettono di scambiarsi mediatori, invece, possono andare ognuno per la sua strada: uno va al bar, l'altro al bagno e l'interazione è spenta.

24.4 Altri processi

Attaccare le linee di campo dei due diagrammi più semplici per mezzo della linea ondulata non è l'unica possibilità. Un altro modo di fare il prodotto $(\phi_1 A \phi_2)^2$ consiste nell'unire due linee di particelle di materia, come in Figura 24.3.

Possiamo interpretare questo diagramma come quello che descrive un campo *libero* (quello a sinistra che proviene dal basso) che interagisce con una particella di materia (a sinistra). In seguito all'assorbimento del campo la particella si propaga per un po' e poi riemette il campo, che si muove verso il basso a destra. Un elettrone che entra in una regione nella quale è presente un campo elettrico non *consuma* il campo presente. Per interagire con esso lo deve assorbire e riemettere. In questo processo cambia la sua direzione perché la cinematica del processo di riemissione può essere diversa da quella dell'assorbimento.

Questo processo è consentito perché possiede almeno due vertici. Si può vedere che il numero di vertici di un diagramma di Feynman coincide con l'or-

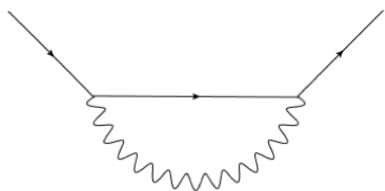


Figura 24.4 L'interazione tra una particella e il suo proprio campo.

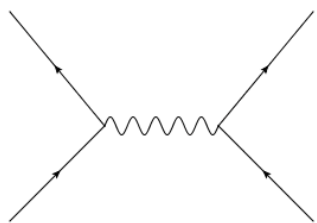


Figura 24.5 Il diagramma di annichilazione.

dine dello sviluppo polinomiale della funzione energia d'interazione. Un altro processo con due vertici è quello che consiste nell'unire due linee di materia e due linee di campo, per ottenere il diagramma di Fig. 24.4.

Questo diagramma rende ragione di quanto accade nel caso in cui la particella in moto emetta un campo e non vi siano particelle nelle vicinanze per *raccoglierlo*. Il campo è semplicemente riassorbito dalla particella.

Ma il diagramma più interessante si ottiene ruotando di 90 gradi quello di Fig. 24.2, mostrato in Fig. 24.5. Anche questo dovrebbe essere un diagramma legittimo, ma che cosa rappresenta? A prima vista sembra un processo impossibile: una particella (un elettrone) proviene da sinistra in basso e raggiunge un vertice nel quale emette un campo. Da questo punto parte anche un'altra particella che però deve muoversi all'indietro nel tempo (finora abbiamo sempre considerato il tempo come se scor-

resse da sinistra a destra). Trascorso qualche istante, poi, il campo svanisce e dal punto in cui svanisce il campo fuoriesce una particella che si muove verso l'alto a destra, ma un'altra particella che si muove all'indietro sembra finire in questo punto e scomparire anche lei! Naturalmente è possibile che abbiamo trascurato qualche particolare secondo il quale questo diagramma dovrebbe restituire il valore zero, una volta calcolato, ma la spiegazione, sorprendente, è nel prossimo paragrafo.

24.5 L'antimateria

Al Paragrafo 16.3 abbiamo detto che nel 1933 fu scoperta una particella identica all'elettrone, ma con carica elettrica positiva: il positrone. Un elettrone in un campo elettrico si muove in modo tale da spostarsi da punti a potenziale minore a punti a potenziale maggiore. Se avesse carica elettrica positiva, il suo moto sarebbe diverso: si muoverebbe spostandosi dai punti a potenziale maggiore a quelli a potenziale minore. Se noi filmassimo un elettrone in un campo elettrico e poi guardassimo il filmato al contrario confonderemmo il moto con quello di un positrone. In altre parole i positroni si comportano come elettroni che si muovono all'indietro nel tempo e viceversa.

La linea di materia presente in Fig. 24.5 a sinistra, dunque, potrebbe benissimo rappresentare un positrone che si muove al contrario rispetto alla direzione della freccia. Quello che ci dice questo diagramma è che quando un elettrone e un positrone si incontrano **annichilano**: vengono distrutti, spariscono nel nulla; la loro materia è completamente annullata, ma la loro energia no. L'energia posseduta si trasferisce al campo che evidentemente è prodotto nell'annichilazione, che si propaga per un po' di tempo, ma successivamente **materializza** in una coppia particella-antiparticella (la linea di materia che si muove al contrario a destra). La cosa interessante è che nei due vertici si devono conservare tutta una serie di grandezze fisiche, ma tra queste non c'è il tipo di particella. Nel vertice di destra dunque si può produrre una coppia elettrone-positrone identi-

ca a quella iniziale (e in questo caso il risultato netto sarebbe un'interazione tra queste due particelle), ma anche una coppia di muoni $\mu^+\mu^-$!

Partendo da una coppia elettrone-positrone si finisce con l'avere una coppia di muoni (o di altre particelle). Evidentemente questo (o uno analogo) è il meccanismo con il quale i raggi cosmici primari danno origine a particelle di natura diversa: nell'urto deve essere emesso un campo che poi materializza in qualche modo. In effetti si può dimostrare sperimentalmente che il fenomeno esiste e funziona proprio come previsto dalla teoria. Ancora una volta, dunque, abbiamo una conferma della bontà della nostra visione del modo in cui procedono le interazioni tra particelle elementari.

24.6 La produzione delle particelle strane

Vediamo come si può interpretare la produzione di particelle strane alla luce di questa teoria. Come illustrato nel Paragrafo 22 i protoni sono particelle costituite di tre quark: $p = (uud)$, mentre i pioni sono composti di una coppia quark-antiquark e quindi $\pi^- = (d\bar{u})$.

Quando queste due particelle si urtano si possono produrre una Λ e uno dei K , che sono composte a loro volta di quark: $\Lambda = (uds)$ e $k = (d\bar{s})$. Nello stato iniziale abbiamo due quark u , due quark d e un quark $\bar{u}$, mentre in quello finale ci sono due quark d , un solo quark u e due quark strani: s e $\bar{s}$. In effetti sembra che ai quark d non accada nulla. Quello che probabilmente succede è che nell'urto si scontrano un quark u del protone con il quark $\bar{u}$ del pione. Trattandosi di una coppia particella-antiparticella possono annichilare emettendo una particella mediatrice di forza. In questo caso molto probabilmente non si produce campo elettromagnetico, ma forte, visto che la sua intensità è molto maggiore. Il mediatore della forza forte è chiamato **gluone**. Il gluone si propaga per un po' e poi materializza in una coppia quark-antiquark diversa: $s\bar{s}$. Questi si legano a quelli che sono stati solo **spettatori** del processo

(non hanno interagito) formando le particelle Λ e K .

24.7 L'interazione debole

Il decadimento dei neutroni procede secondo la reazione $n \rightarrow p + e^- + \nu$ mediato dall'**interazione debole**. Quello che deve accadere è che un quark d del neutrone si trasforma in un quark u formando il protone, con la conseguente emissione di un elettrone e un neutrino. L'interazione debole non *vede* i diversi tipi di quark: il quark d e il quark u sono diversi solo per l'interazione forte e per quella elettromagnetica (avendo carica diversa). L'interazione debole praticamente non distingue tra questi due quark, quindi il quark d può diventare un quark u emettendo un campo debole. Ma dal momento che in questo processo cambia la carica elettrica del quark, il campo debole deve trasportare esso stesso una carica elettrica negativa. Questo campo, mediato dai **bosoni** W , elettricamente carichi, produce una coppia di particelle che nuovamente sono identiche dal punto di vista dell'interazione debole (o meglio sono l'una l'antiparticella dell'altra), ma non per quella elettromagnetica: un elettrone e un neutrino.

Le interazioni deboli possono essere mediate anche da un'altra particella: la **Z**, che è neutra. La Z è responsabile degli eventi nei quali i neutrini urtano un protone e gli trasferiscono una quantità significativa di energia nel processo $\nu + p \rightarrow \nu + p$. Sperimentalmente questo processo fu osservato per la prima volta al CERN nel 1973. Nell'esperimento si inviava un fascio di neutrini su un bersaglio (costituito di protoni e neutroni). Quando un neutrino urta un protone lo fa muovere nel rivelatore provocando una traccia ionizzata, mentre il neutrino prosegue la sua corsa, dopo essere stato deviato, senza lasciare tracce nel rivelatore.

Il bosone di Higgs

FILMATO NON RIPRODUCIBILE SU QUESTO
SUPPORTO: DIGITA L'URL NELLA CAPTION O
SCARICA L'E-BOOK

Figura 25.1 La massa di una particella è una misura di quanto sia difficile cambiarne lo stato di moto. Da questo punto di vista il meccanismo di Higgs è analogo all'effetto prodotto da un campo magnetico su una biglia d'acciaio [<https://www.youtube.com/watch?v=Dkd0--yxI0w>].

La scoperta del bosone di Higgs avvenuta nel 2012 a opera degli esperimenti ATLAS e CMS all'acceleratore LHC del CERN rappresenta una delle imprese scientifiche più ardite che l'uomo abbia mai realizzato. Per renderla possibile sono infatti state impiegate tecniche al limite della tecnologia.

Dopo la scoperta, che ha avuto grande enfasi su tutti i media, è risultato abbastanza noto a tutti che il bosone di Higgs è la particella responsabile del fatto che tutte le altre particelle possiedono una massa. Tuttavia, essendo abituati a pensare alla massa come a una proprietà intrinseca della materia, risulta difficile immaginare perché ci sia bisogno di un meccanismo per dare massa alle particelle e come sia possibile che questa proprietà emerga dall'interazione con un'altra particella.

In questo capitolo descriviamo quello che si chiama il *meccanismo di Higgs*, che **Peter Higgs** teorizzò nel 1964 per spiegare un complesso problema di fisica fondamentale che ha a che fare con l'osservazione sperimentale della rottura di alcune simmetrie dell'Universo. La spiegazione originale di Higgs ser-

viva di fatto a spiegare la differenza di massa tra i fotoni, responsabili dell'interazione elettromagnetica, e i bosoni vettori intermedi Z e W , responsabili dell'**interazione debole**, attraverso l'introduzione di una nuova particella, successivamente battezzata **bosone di Higgs**. È quindi possibile estendere in modo abbastanza naturale le interazioni di questa particella per fare in modo che questa possa dare la massa anche alle particelle di materia (quark e leptoni).

La spiegazione di Higgs richiede l'introduzione della teoria quantistica dei campi, che va molto al di là degli scopi di questa pubblicazione, ma si può riformulare [32] in termini di fisica classica rovesciando il problema: spiegando cioè prima l'acquisizione della massa da parte delle particelle di materia e poi la differenza di massa tra i fotoni e i bosoni vettori intermedi (Z e W).

25.1 Richiami sul concetto di energia

Il concetto di energia è uno dei più ostici per gli studenti, nonostante il fatto che, tutto sommato, il calcolo dell'energia di un corpo in una determinata condizione sia relativamente semplice. In questo contesto c'interessa osservare come, a dispetto delle apparenze, il calcolo dell'energia di un corpo in condizioni molto diverse, sia sempre esprimibile nella stessa maniera. Nel seguito consideriamo sempre corpi che, nel sistema di riferimento scelto, sono in quiete, per cui la loro energia cinetica è nulla.

Consideriamo inizialmente un corpo di massa m in un campo gravitazionale $\mathbf{G}$. L'energia assunta

dal corpo in virtù dell'interazione con il campo è esprimibile come

$$U_G = m\mathcal{G}, \quad (25.1)$$

dove $\mathcal{G}$ è il cosiddetto *potenziale gravitazionale*. Vale la pena ricordare che il potenziale di un campo è una funzione scalare del campo stesso e delle coordinate. Nello specifico, scegliendo un punto a distanza infinita come riferimento, abbiamo che

$$\mathcal{G} = \int_{\infty}^r \mathbf{G} \cdot d\mathbf{r}. \quad (25.2)$$

Nell'equazione (25.2), $\mathbf{r}$ rappresenta il vettore che individua la posizione del corpo nel sistema di riferimento scelto. Nel caso semplice in cui il campo gravitazionale sia prodotto da un corpo di massa M possiamo dunque scrivere che

$$U_G = GM \int_{\infty}^r \frac{\mathbf{r} \cdot d\mathbf{r}}{r^3}, \quad (25.3)$$

dove G è la costante di Newton. Data l'arbitrarietà con la quale si può scegliere il punto nel quale $U_G = 0$, l'energia del corpo è definita a meno di una costante che, proprio per quanto sopra, possiamo sempre scegliere essere uguale a zero (questa scelta è d'ora in poi considerata implicita).

Consideriamo ora un corpo con carica elettrica q immerso in un campo elettrostatico $\mathbf{E}$. Anche per il campo elettrostatico possiamo definire un potenziale $\mathcal{E}$ ¹ in maniera del tutto analoga a quanto fatto per il campo gravitazionale e in definitiva scrivere che

$$U_E = q\mathcal{E}. \quad (25.4)$$

Il potenziale $\mathcal{E}$ è ancora una volta una funzione scalare del campo e delle coordinate, la cui forma è identica a quella dell'equazione (25.2), avendo cura di sostituire $\mathbf{E}$ a $\mathbf{G}$.

¹Di solito il potenziale elettrostatico si indica col simbolo V o ΔV , ma in questo caso preferiamo adoperare il simbolo $\mathcal{E}$ per rendere evidente il tipo di campo a cui si riferisce e per evitare di confondere il potenziale con il volume, che indichiamo, invece, con V .

Nel caso di una spira di area S , percorsa da corrente I e immersa in un campo magnetico $\mathbf{B}$, l'energia assunta dalla spira vale

$$U_B = -IS\hat{\mathbf{z}} \cdot \mathbf{B} \quad (25.5)$$

dove $\hat{\mathbf{z}}$ è un versore orientato in modo da essere perpendicolare al piano su cui si giace la spira. Solitamente la quantità $\mathbf{m} = IS\hat{\mathbf{z}}$ è chiamata *momento magnetico* della spira $\boldsymbol{\mu}$, per cui si scrive che $U_B = -\boldsymbol{\mu} \cdot \mathbf{B}$. È ben noto che, nel caso del campo magnetico, non è possibile scrivere un potenziale scalare, ma abusando leggermente del vocabolario possiamo ridefinire il termine *potenziale* come un'opportuna funzione scalare dei campi e delle coordinate, tale per cui possiamo scrivere che

$$U_B = I\mathcal{B}. \quad (25.6)$$

Dal confronto delle ultime due equazioni si deduce immediatamente che il *potenziale* (così come da noi ridefinito) $\mathcal{B}$ del campo magnetico vale

$$\mathcal{B} = -S\hat{\mathbf{z}} \cdot \mathbf{B}. \quad (25.7)$$

Sarebbe forse più opportuno definire un termine *ad hoc* per questa funzione, invece di usare il termine *potenziale*, ma per semplicità continuiamo a impiegare questo nome, scritto in caratteri diversi. In definitiva si può osservare come l'energia di un corpo immerso in un campo si possa scrivere sempre nella stessa forma, per tutti i campi noti a uno studente di liceo:

$$U = U_G + U_E + U_B = m\mathcal{G} + q\mathcal{E} + I\mathcal{B}. \quad (25.8)$$

Vale la pena osservare che la sorgente del campo gravitazionale è la massa, quella del campo elettrostatico la carica elettrica e quella del campo magnetico la corrente. Tutti i termini di questa somma hanno la stessa forma: una costante di accoppiamento che dipende dalla natura del campo con il quale la particella in esame interagisce (m , q o I) e il *potenziale* del campo di cui la particella stessa è sorgente ($\mathcal{G}$, $\mathcal{E}$ o $\mathcal{B}$).

25.2 Campi autointeragenti

Quanto sopra tiene conto dell'energia posseduta dai corpi immersi nei campi. È però noto anche a studenti di liceo che i campi elettrici e magnetici trasportano energia. Basta considerare un condensatore carico per rendersi conto che l'energia in esso contenuta non può che essere trasportata dal campo elettrico tra le armature; analizzando un circuito RLC, invece, si vede subito che il campo magnetico nell'induttanza deve poter trasportare una certa quantità di energia.

Nel caso classico l'energia del campo è distribuita all'interno di un volume (quello tra le armature del condensatore o all'interno della bobina negli esempi sopra riportati) e si parla dunque di densità di energia dei campi. Per i campi elettrici e magnetici nel vuoto le densità di energia u_E^a e u_B^a si calcolano rispettivamente come

$$u_E^a = \frac{\varepsilon_0}{2} \mathbf{E} \cdot \mathbf{E} \quad (25.9)$$

e

$$u_B^a = \frac{1}{2\mu_0} \mathbf{B} \cdot \mathbf{B}. \quad (25.10)$$

Nelle equazioni sopra riportate l'indice a sta a indicare il fatto che questi contributi all'energia derivano dall'*autointerazione* dei campi in questione con sé stessi. Questo spiega perché un tale termine non sia presente per il campo gravitazionale, che non interagisce con sé stesso (possiamo comunque pensare che la sua densità di energia sia $u_G^a = \frac{\gamma}{2} \mathbf{G} \cdot \mathbf{G}$ con $\gamma = 0$). In un volume V in cui siano presenti sia campi elettrici che magnetici, l'energia contenuta dovuta alla presenza di questi campi è quindi

$$U_E^a + U_B^a = V \left(\frac{\varepsilon_0}{2} \mathbf{E} \cdot \mathbf{E} + \frac{1}{2\mu_0} \mathbf{B} \cdot \mathbf{B} \right). \quad (25.11)$$

25.3 Sul significato dell'energia

Considerato quanto sopra, l'energia contenuta in un volume V di Universo, nel quale sia presente una

particella di massa m e carica elettrica q , che dunque generi una corrente $I = dq/dt$ se in moto, in presenza di campi elettrici e magnetici, si può scrivere come

$$U = \frac{1}{2} m \mathbf{v} \cdot \mathbf{v} + m\mathcal{G} + q\mathcal{E} + I\mathcal{B} + V \left(\frac{\varepsilon_0}{2} \mathbf{E} \cdot \mathbf{E} + \frac{1}{2\mu_0} \mathbf{B} \cdot \mathbf{B} \right). \quad (25.12)$$

A parte il termine cinetico (che peraltro ha una forma analoga a quelli dovuti all'autointerazione dei campi suggerendo che questi ultimi dovrebbero avere un'analogia natura, come si scopre con la meccanica quantistica dei campi), tutti gli altri termini sono il risultato di un'interazione. Questo suggerisce che, se potessimo *spegner*e tutte le interazioni dell'Universo, l'energia contenuta in quest'ultimo sarebbe nulla (almeno quella potenziale; su questo argomento si potrebbero fare alcune considerazioni che tuttavia esulano dallo scopo di questo articolo e per semplicità, per il momento, ci occuperemo soltanto dei termini non cinetici). Di fatto possiamo interpretare l'energia come una grandezza fisica che caratterizza l'intero Universo il cui valore deve rimanere costante nel tempo. Una variazione dell'energia in una regione dell'Universo comporta una variazione contraria in una regione diversa e questo equivale al manifestarsi di un'interazione. Infatti, secondo l'equazione (25.12) si può avere una variazione dell'energia in una regione solo a seguito dall'accensione di un'interazione oppure della modifica della velocità della particella (il che però implica un'accelerazione, dunque una forza, dunque un'interazione).

In definitiva potremmo dire che in un Universo privo d'interazioni l'energia sarebbe nulla o meglio che l'energia contenuta in una regione qualsiasi di Universo dipende esclusivamente dalla presenza dei campi e più precisamente dal fatto che le particelle interagiscono con questi oppure dal fatto che i campi possono, in certi casi, interagire con sé stessi.

Tenendo conto di ciò possiamo esprimere l'energia contenuta in una regione di volume V nella quale siano presenti campi e particelle in quiete come

$$U = \sum_i \alpha_i \mathcal{F}_i + V \sum_i \beta_i \mathbf{F}_i \cdot \mathbf{F}_i, \quad (25.13)$$

dove i coefficienti α_i rappresentano le costanti di accoppiamento tra particelle e campi (che dipendono dalle caratteristiche delle particelle che sono a loro volta sorgenti dello stesso campo) e $\mathcal{F}_i$ opportune funzioni dei campi $\mathbf{F}_i$ che abbiamo chiamato *potenziali*. I coefficienti β_i invece rappresentano le costanti di accoppiamento dei campi con sé stessi.

Ad esempio, consideriamo un condensatore con armature di superficie S distanti d con il vuoto come dielettrico, con all'interno una particella di carica elettrica q vincolata in un punto a distanza δ dall'armatura a potenziale più basso. L'energia in esso contenuta vale

$$U = q\mathcal{E}(\delta) + Sd\frac{\varepsilon_0}{2}E^2 \quad (25.14)$$

con $\mathcal{E}(\delta) = E\delta$ ed $E = |\mathbf{E}|$, trascurando la gravità. Confrontando quest'espressione con la (25.13) vediamo che c'è un solo termine ($i = 1$) per ciascuno dei due addendi per cui $\alpha_1 = q$, $\mathcal{F}_1 = \mathcal{E}(\delta)$, $\beta_1 = \varepsilon_0/2$ e $\mathbf{F}_1 = \mathbf{E}$. Se $\mathbf{E}$ fosse nullo $U = 0$.

25.4 L'introduzione della relatività

Quanto detto finora sul significato fisico dell'energia comincia a vacillare non appena si tenga conto della relatività speciale. In questo caso è noto che all'energia cinetica e a quella dovuta alle interazioni deve essere aggiunto un termine mc^2 detto, per l'appunto, energia a riposo della particella.

Con la relatività viene a cadere il principio (del tutto arbitrario, se vogliamo, ma ragionevole) secondo il quale l'energia contenuta in una regione di spazio dipenda esclusivamente dall'interazione di qualcosa con qualcos'altro. Il termine mc^2 non ha affatto la forma degli altri termini. Dipende esclusivamente dalla natura della particella considerata e sarebbe non nullo anche in assenza di ogni interazione.

È abbastanza naturale chiedersi (sebbene nessuno se lo sia chiesto per molti decenni) perché mai un tale termine debba entrare nella determinazione dell'energia di una particella quando tutti gli altri dipendono dal fatto che la particella in questione interagisce con un campo. E ammettendo che questo sia del tutto ragionevole, perché non esiste un termine del tipo qk^2 dove k^2 sia una combinazione di costanti aventi le dimensioni di un'energia per unità di carica elettrica?

25.5 Il Meccanismo di Higgs

A queste domande risponde il Meccanismo di Higgs. Supponiamo di tornare nella condizione prevista dalla fisica classica, secondo la quale vale l'equazione (25.13). Supponiamo, inoltre, che oltre ai campi già noti ne esista un altro che indichiamo con W , il cui *potenziale* indicheremo con $\mathcal{W}$. Si noti che il campo W è un campo scalare e non un campo vettoriale. Per quanto stiamo considerando, tuttavia, questo non fa differenza.

In presenza di questo nuovo campo l'equazione (25.13) diventa

$$U = \sum_i \alpha_i \mathcal{F}_i + a\mathcal{W} + V \left(\sum_i \beta_i \mathbf{F}_i \cdot \mathbf{F}_i + bW \cdot W \right), \quad (25.15)$$

dove nelle somme abbiamo inclusi i soli campi già noti per rendere esplicito il fatto che il campo W non è tra questi. Il campo è peculiare per un'altra ragione: a differenza di tutti gli altri campi, il cui potenziale $\mathcal{F}_i$ assume il valore minimo quando il campo è nullo, il campo W possiede un potenziale il cui minimo si ottiene per $W = W_0 \neq 0$. Vedremo più avanti il significato di questa scelta. Per il momento prendiamola per buona e scriviamo il campo $W = W_0 + H$. Corrispondentemente $\mathcal{W} = \mathcal{W}_0 + \mathcal{H}$. L'energia si scrive dunque come

$$U = \sum_i \alpha_i \mathcal{F}_i + a(\mathcal{W}_0 + \mathcal{H}) + V \left(\sum_i \beta_i \mathbf{F}_i \cdot \mathbf{F}_i + b(W_0 + H)^2 \right) \quad (25.16)$$

ed espandendo il quadrato otteniamo

$$U = \sum_i \alpha_i \mathcal{F}_i + a\mathcal{W}_0 + a\mathcal{H} + V \left(\sum_i \beta_i \mathbf{F}_i \cdot \mathbf{F}_i + bW_0^2 + bH^2 + 2bW_0H \right). \quad (25.17)$$

Ora consideriamo uno per uno i termini in piú comparsi nell'espressione rispetto a quella originale.

Il termine $a\mathcal{W}_0$ è una costante che dipende unicamente dalle caratteristiche della particella presente nella regione di spazio considerata. a , infatti, è la costante di accoppiamento di questa particella al campo minimo il cui potenziale è $\mathcal{W}_0$. In altre parole, questo termine è del tutto analogo a $m\mathcal{G}$ o a $q\mathcal{E}$. L'unica differenza è che il valore del campo con il quale la particella interagisce è costante in tutto l'Universo e vale W_0 . Di fatto $a\mathcal{W}_0$ rappresenta l'energia d'interazione di una particella con un campo la cui intensità è costante. Questo termine perciò dev'essere uguale dappertutto e può solo dipendere dalla natura della particella attraverso la costante d'accoppiamento a . Se per accidente $a\mathcal{W}_0 = mc^2$ questo termine rappresenta proprio l'energia a riposo di una particella.

È così che una particella priva di massa con *carica di Higgs* a acquista una massa m interagendo con il campo nella sua configurazione di minima energia W_0 . Ma cosa vogliono dire gli altri termini?

Il prodotto $a\mathcal{H}$ rappresenta l'interazione delle particelle con il campo di Higgs in eccedenza rispetto al valore che rende minimo il suo potenziale. In altre parole è possibile che in una regione di spazio sia presente un campo di Higgs piú intenso rispetto a quello nel quale il potenziale di Higgs è minimo. Questo non conduce a una maggiore o minore massa per la particella, ma a un'interazione tra la particella e il campo residuo per certi versi paragonabile a quella di una carica con un campo elettrico. Possiamo pensare a questo tipo d'interazione come a un'attrazione o a una repulsione della particella da parte del campo. Il termine $a\mathcal{H}$ infatti non dipende solo dalla particella, ma è del tutto analogo al termi-

ne $q\mathcal{E}$ e pertanto dipende dalla particella (attraverso a) e dal campo (attraverso $\mathcal{H}$).

bW_0^2 è un termine del tutto irrilevante ai fini della dinamica. Infatti questo è davvero un termine costante, che dipende solamente dal campo nella sua configurazione di minima energia e dalla sua interazione con sé stesso. Questo addendo è uguale in tutti i punti dell'Universo e può essere eliminato ridefinendo la costante additiva arbitraria come $-bW_0^2$.

Il termine bH^2 , analogo a $\varepsilon_0 E^2/2$, rappresenta l'interazione del campo di Higgs in eccesso con sé stesso, mentre quello che si scrive come $2bW_0H$ rappresenta l'interazione del campo H con il campo W_0 . Da una parte questo è analogo a $\varepsilon_0 E^2/2$ rappresentando l'interazione di un campo con un'altro della sua stessa natura; dall'altra il termine in questione è simile, per certi versi, a $a\mathcal{W}_0$, perché uno dei campi coinvolti è proprio quello che determina il valore minimo del potenziale. In sostanza questo addendo nell'energia dipende solo dal campo di Higgs in eccesso (perché tutto il resto è costante) e cresce proporzionalmente alla *quantità* di campo presente. Rappresenta dunque quello che potremmo chiamare la *massa del campo*. Il campo di Higgs, dunque, è un campo massivo.

25.6 Sulla forma del potenziale di Higgs

Nella sezione precedente abbiamo visto che il campo di Higgs possiede due caratteristiche che, in qualche modo, lo rendono diverso dai campi elettrico e magnetico:

- il valore minimo del potenziale di autointerazione non si ottiene per $W = 0$, ma per $W \neq 0$.
- è scalare.

Un modo per far sí che l'energia sia minima quando il campo è diverso da zero è il seguente. Supponiamo di avere un campo autointeragente W . Seguendo

l'analogia finora esposta con i campi elettrici e magnetici scriveremmo il contributo all'energia dovuto all'autointerazione come

$$U = \alpha W^2 \quad (25.18)$$

ma così facendo per $W = 0$ si avrebbe $U = 0$ che corrisponde al valore minimo dell'energia. Si vede subito, infatti, che U è una parabola nel piano (W, U) . Senza rinunciare all'ipotesi secondo la quale i contributi all'energia sono imputabili alle interazioni tra particella e campo e tra campo e campo, possiamo tranquillamente assumere che l'interazione tra campi di Higgs proceda in modo tale da contribuire all'energia secondo l'equazione

$$U = \alpha W^2 + \beta W^4. \quad (25.19)$$

In fondo, il termine W^2 rappresenta il solito contributo all'energia dovuto all'autointerazione dei campi, ma dal momento che W è un campo scalare, il termine W^2 si può pensare anch'esso come un campo scalare che autointeragendo dà luogo al termine βW^4 . Nel caso dei campi vettoriali questo non avviene perché il campo ha carattere vettoriale, mentre l'energia è uno scalare. Ma nel caso di campi scalari è possibile. Seguendo questa linea di pensiero potremmo anche ammettere l'esistenza di termini con potenze superiori di W . Questo non è escluso, ma è ragionevole pensare che la probabilità d'interazione diminuisca fortemente all'aumentare del numero di campi per cui possiamo trascurare i termini di ordine superiore. D'altra parte questo potrebbe anche voler dire che l'espressione sopra ricavata non è altro che l'espansione in serie di Taylor di qualche funzione più complessa del campo W .

Troviamo il minimo di (25.19): questo si ha quando

$$\frac{dU}{dW} = 2\alpha W + 4\beta W^3 = 0, \quad (25.20)$$

e cioè per

$$W^2 = -\frac{\alpha}{2\beta}. \quad (25.21)$$

Se β e α sono discordi (prendiamo, tanto per fissare le idee $\alpha < 0$ e $\beta > 0$) il minimo del potenziale

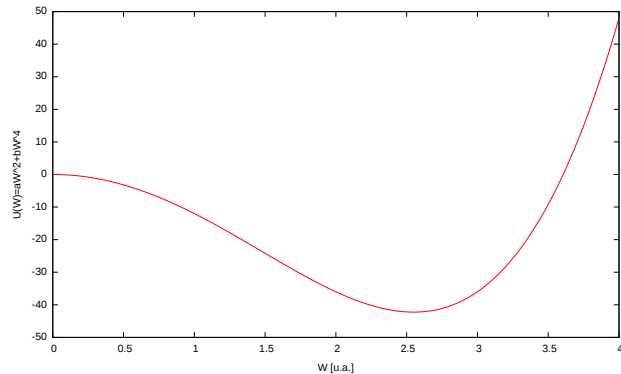


Figura 25.2 Il potenziale del campo di Higgs in funzione dell'intensità del campo in unità arbitrarie. Abbiamo scelto $a = -13$ e $b = 1$.

si ha naturalmente per $W \neq 0$. Se si fa un grafico dell'andamento di U in funzione di W si ottiene la figura mostrata in Fig. 25.2.

Dalla figura si vede che il minimo dell'energia si ottiene quando $W = \sqrt{-a/2\beta} \simeq 2.5$, avendo scelto $a = -13$ e $b = 1$. L'energia che si ottiene in questo modo è negativa, ma è sufficiente scegliere una costante arbitraria opportuna da sommare a questo valore per renderla positiva o nulla. Lo stato di vuoto *classico* (quello in cui particelle e campi sono assenti) possiede dunque un'energia maggiore di uno stato in cui è presente una certa quantità di campo. Questo implica che un tale stato di vuoto è instabile e tende spontaneamente a evolvere in uno stato di energia più bassa, nel quale è presente un campo di Higgs non nullo. Dobbiamo quindi pensare che il *vero* stato di vuoto sia quello nel quale abbiamo rimosso tutti i campi e le particelle, tranne il campo di Higgs che si riformerà spontaneamente anche se riuscissimo a rimuoverlo. Per questo lo stato nel quale $W = W_0$ si può considerare l'effettivo stato di vuoto che non è quello in cui *non c'è nulla*, ma quello di minima energia.

Questo spiega anche perché abbiamo scelto un campo scalare. Un primo argomento è, come abbiamo visto, che il campo scalare può interagire con sé stesso dando luogo a contributi all'energia con potenze maggiori. Inoltre il campo di Higgs quando

si trova nello stato di minima energia deve essere rappresentativo dello stato di vuoto (che è quello stato nel quale non si misura nulla). Il vuoto non può avere una direzione privilegiata come avrebbe se il campo di Higgs fosse vettoriale.

Possiamo farci un'idea abbastanza precisa di quel che accade con un'altra analogia. Consideriamo una piscina di palline, come quelle che si trovano nei parchi giochi o in alcuni centri commerciali per intrattenere i bambini mentre i genitori fanno shopping. Ammettiamo che la nostra piscina di palline rappresenti un volume di Universo che possiamo osservare. Se la osserviamo da fuori, seduti su una panchina, le palline più in superficie sono sotto il bordo e non sono visibili. La piscina, per noi, è vuota. Non nel senso usuale del termine (la piscina è piena di palline), ma nel senso che è impossibile osservare più di quanto riusciamo a vedere in queste condizioni. Un bambino fuori dalla piscina si muove liberamente, come una particella a massa nulla. Ma se lo facciamo entrare nella regione di spazio in cui è presente il campo di Higgs rappresentato dalle palline, si muoverà con difficoltà. Se si muove lentamente non smuoverà le palline abbastanza da renderle visibili e noi potremo concludere che la massa del bambino nella piscina è maggiore perché occorre una forza maggiore per accelerarlo. Se si muovesse più rapidamente la sua interazione con il campo aumenterebbe e potrebbe causare la comparsa di un campo misurabile (vedremmo saltellare di quando in quando delle palline oltre il bordo). Per noi è come osservare un aumento del campo. Questo fenomeno è quello descritto dal termine $a\mathcal{H}$ dell'energia. Se il campo in eccesso è abbastanza intenso, poi, potremmo osservare anche l'interazione del campo residuo con sé stesso bH^2 come palline che si toccano. Il termine bW_0^2 invece rappresenta l'energia delle palline sotto il bordo: perché sia possibile che le palline arrivino abbastanza vicine al bordo da poter essere osservate in presenza di altre particelle è necessario che le palline interagiscano tra loro: e in effetti lo fanno perché quando una sta sopra l'altra quella in alto non può stare più in basso di quanto sia. Infine il termine $2bW_0H$ rappresenta l'interazione tra campo residuo e campo minimo che si può raffigurare come

quella che esiste tra una pallina in volo (visibile) che rimbalza sulle palline sotto il bordo della piscina più in superficie.

25.7 Campi massivi

L'idea di un campo con massa è forse la più difficile da assimilare per uno studente. Per comprendere cosa sia un campo con massa possiamo fare così: supponiamo di avere una particella di massa M che produce un campo di qualche tipo (per esempio gravitazionale). Il campo prodotto è privo di massa e per questo la particella conserva la sua massa M . Se il campo prodotto avesse massa m , per la conservazione di quest'ultima, la massa della particella che lo produce dovrebbe diminuire e diventare $M - m$. Ma così facendo, prima o poi la particella perderebbe tutta la sua massa. Se infatti in un punto P c'è un campo $\mathbf{G}$, trascorso un certo tempo Δt , questo campo si è propagato in un punto che dista $c\Delta t$ da P , dove c è la velocità di propagazione del campo (nel caso del campo elettrico sarebbe la velocità della luce). Per fare in modo che in P continui a esserci un campo $\mathbf{G}$ la sorgente deve rimpiazzarlo perdendo un'ulteriore frazione della sua massa e così via.

Se ammettiamo che la particella possa perdere una frazione della sua massa per un tempo limitato Δt_{max} possiamo però ammettere che produca un campo massivo di massa m (diventando una particella di massa $M - m$) il quale, trascorso questo tempo, non può più esistere e deve essere per così dire riassorbito dalla particella che lo ha prodotto. La particella sorgente infatti deve tornare ad avere la massa originale M trascorso il tempo Δt_{max} . Questo implica che il campo può al massimo raggiungere una distanza dalla particella pari a circa

$$L_{max} \simeq c\Delta t_{max}. \quad (25.22)$$

In altre parole un campo massivo non è altro che un campo a raggio limitato. Oltre una certa distanza dalla sorgente, non si osserva più alcun campo.

25.8 La massa dei bosoni vettori

Riscriviamo l'equazione che ci fornisce l'energia contenuta in un volume V di Universo in un caso particolare: prendiamo come volume V quello all'interno di un condensatore carico all'interno del quale sia presente una particella elettricamente carica, con carica q . Trascurando la gravità, l'energia contenuta in questo condensatore è la somma dell'energia elettrostatica della carica q e di quella immagazzinata sotto forma di campo elettrico. Se non avessimo il campo di Higgs quest'energia ammonterebbe a

$$U = U_{em} = q\mathcal{E} + V\frac{\epsilon_0}{2}\mathbf{E} \cdot \mathbf{E}. \quad (25.23)$$

In presenza anche di **interazioni deboli**, all'energia dovremmo sommare un termine del tipo

$$U_{weak} = w\mathcal{Z} + V\frac{\zeta_0}{2}\mathbf{Z} \cdot \mathbf{Z} \quad (25.24)$$

dove $\mathbf{Z}$ rappresenta il campo debole, $\mathcal{Z}$ il suo *potenziale* e ζ_0 è una costante che si deve determinare sperimentalmente e che indica l'intensità dell'autointerazione del campo debole. La simmetria tra le due espressioni è evidente, perciò possiamo riscrivere tutto come

$$U = E_{em} + U_{weak} = \mathbf{c} \cdot \mathcal{D} + V\frac{1}{2}\mathbf{D} \cdot \mathbf{D}, \quad (25.25)$$

dove $\mathbf{D}$, $\mathcal{D}$ e $\mathbf{c}$ sono vettori di due componenti. Le componenti del vettore $\mathbf{D}$ sono i moduli di due vettori spaziali:

$$\mathbf{D} = \begin{pmatrix} \sqrt{\epsilon_0} E \\ \sqrt{\zeta_0} \mathcal{Z} \end{pmatrix}. \quad (25.26)$$

Le componenti di $\mathcal{D}$ sono

$$\mathcal{D} = \begin{pmatrix} \mathcal{E} \\ \mathcal{Z} \end{pmatrix}, \quad (25.27)$$

mentre quelle del vettore $\mathbf{c}$ sono le costanti di accoppiamento ai rispettivi campi:

$$\mathbf{c} = \begin{pmatrix} q \\ w \end{pmatrix}. \quad (25.28)$$

Questi vettori, naturalmente, non sono vettori nello spazio ordinario: *vivono* in uno spazio astratto bi-dimensionale i cui assi sono allineati lungo direzioni che dipendono dalle interazioni. In altre parole, la direzione di questi vettori nello spazio astratto definisce in qualche maniera l'intensità relativa tra le diverse interazioni. In un Universo in cui ci siano solo interazioni elettromagnetiche, il vettore $\mathbf{D}$ assumerebbe una data direzione in questo spazio astratto, mentre in un Universo in cui siano presenti sia interazioni elettromagnetiche che deboli, il vettore formerebbe un angolo non nullo con quello precedente. Infine, in un Universo in cui le interazioni elettromagnetiche scomparissero, il vettore $\mathbf{D}$ avrebbe una direzione perpendicolare al primo.

Ripetendo il ragionamento fatto sopra circa la necessità d'introdurre un nuovo campo per giustificare la presenza di un termine di massa nell'espressione dell'energia, dovremmo aggiungere all'equazione (25.25) un campo ϕ che, per potersi sommare a quelli già presenti, deve essere rappresentato da un vettore di campi:

$$\phi = \begin{pmatrix} \phi_1 \\ \phi_2 \end{pmatrix}. \quad (25.29)$$

Il campo ϕ deve essere autointeragente e presentare un termine di autointerazione del tipo ϕ^4 . In questo caso l'energia si scriverebbe come

$$U = \mathbf{c} \cdot \mathcal{D} + V\frac{1}{2}\mathbf{D} \cdot \mathbf{D} + \mathbf{a} \cdot \Phi + Vb\phi \cdot \phi + Vc\phi^4 + g\mathbf{V} \cdot \phi, \quad (25.30)$$

dove Φ rappresenta il *potenziale* di ϕ . In questo spazio il potenziale del campo ha la forma che si ottiene facendo ruotare la curva della Figura 25.2 attorno all'asse delle ordinate: sarebbe quindi una superficie con la forma di un *sombrero*. È evidente che in questo caso non c'è un solo minimo del potenziale, ma ce ne sono infiniti. Ogni stato di minimo è equivalente all'altro e rappresenta uno stato di minima

energia del tutto simmetrico a ogni altro che possiamo scegliere. Se scegliamo uno dei possibili stati di minima energia, fissiamo le coordinate (ϕ_1^0, ϕ_2^0) in questo spazio e possiamo scrivere il vettore ϕ come

$$\phi = \begin{pmatrix} \phi_1^0 + \eta_1 \\ \phi_2^0 + \eta_2 \end{pmatrix}. \quad (25.31)$$

In questo modo l'interazione del campo $\mathbf{V}$ col campo ϕ , $g\mathbf{V} \cdot \phi$, fa apparire due addendi nell'energia:

$$gE\sqrt{\varepsilon_0}(\phi_1^0 + \eta_1) + gZ\sqrt{\zeta_0}(\phi_2^0 + \eta_2). \quad (25.32)$$

A questo punto osserviamo che $gE\sqrt{\varepsilon_0}\phi_1^0$ e $gZ\sqrt{\zeta_0}\phi_2^0$ sono costanti che possono essere pensate come a quei termini che danno origine ai contributi delle masse dei campi $\mathbf{E}$ e $\mathbf{Z}$. La massa del campo $\mathbf{E}$, come sappiamo, è nulla, mentre quella del campo $\mathbf{Z}$ non lo è. Essendo però tutti gli stati di minima energia del potenziale di Higgs equivalenti tra loro, potremmo certamente scegliere uno stato per cui $\phi_1^0 = 0$ e si avrebbe che $\phi = (0, \phi_2)$. Non abbiamo nessuna ragione per preferire uno stato di minimo piuttosto che un altro e l'equazione che ci dà l'energia è perfettamente simmetrica, ma una volta scelto lo stato di minimo la simmetria viene **rotta**. Del resto, anche la direzione degli assi che definiscono le diverse interazioni sono in un certo senso arbitrarie. Quello che possiamo pensare, perciò, è che la direzione dell'asse delle interazioni elettromagnetiche sia proprio quella definita dal punto di minima energia del campo di Higgs, mentre quella dell'asse delle interazioni deboli sia quello perpendicolare. In questo modo i mediatori delle interazioni elettromagnetiche vengono ad assumere automaticamente una massa nulla, mentre quelli delle interazioni deboli acquistano una massa.

In definitiva possiamo ritenere che prima dell'apparizione dell'Universo non fosse presente alcun campo né materia. Alla nascita dell'Universo compare una certa quantità di campo di Higgs, la cui energia non è necessariamente la minima possibile. Di conseguenza i campi di Higgs cominciano a interagire portando l'energia dal valore assunto alla

nascita al valore minimo possibile. Di stati nei quali avviene questo ce ne sono infiniti e naturalmente l'Universo *cadra* in uno solo di questi infiniti stati. Questo definisce una direzione privilegiata nello spazio astratto delle interazioni per cui lungo questa direzione si sviluppa un tipo d'interazione a raggio infinito e a massa nulla, mentre nelle direzioni perpendicolari se ne sviluppano altre a corto raggio.

Fisicast

Il bosone di Higgs è stato oggetto di una puntata speciale di Fisicast: <http://www.radioscienza.it/2012/07/24/il-bosone-di-higgs/>.

Appendice

In questa sezione sono illustrate alcune tecniche generali, per lo piú di carattere matematico, utili per la soluzione di piú problemi ed esercizi.

Approssimazione di funzioni

In certi casi alcune funzioni complicate si possono approssimare con polinomi di grado e coefficienti opportuni. In altre parole

$$f(x) \simeq \sum_{i=0}^n a_i x^i. \quad (25.33)$$

In molti casi è sufficiente fermarsi a $n = 1$, cioè approssimare una funzione con una retta. In questo paragrafo vediamo alcuni esempi.

La funzione radice quadrata

Se x è piccolo la funzione $f(x) = \sqrt{1 + \alpha x}$ si può approssimare con la retta $p(x) = a_0 + a_1 x$. Troviamo i valori di a_0 e di a_1 .

Per $x = 0$, $f(x) = 1$, mentre $p(0) = a_0$. Impo-
nendo che $f(0) = p(0)$ si trova che $a_0 = 1$. Quando
invece $x \neq 0$ abbiamo

$$\sqrt{1 + \alpha x} \simeq 1 + a_1 x \quad (25.34)$$

per cui, elevando al quadrato

$$1 + \alpha x \simeq 1 + a_1^2 x^2 + 2a_1 x. \quad (25.35)$$

Se x è piccolo, il termine $a_1^2 x^2$ è trascurabile e

$$1 + \alpha x \simeq 1 + 2a_1 x \quad (25.36)$$

e quindi $a_1 = \frac{\alpha}{2}$. In definitiva

$$\sqrt{1 + \alpha x} \simeq 1 + \frac{\alpha}{2} x. \quad (25.37)$$

Logaritmo ed esponenziale

Scriviamo il logaritmo di $1 + x$ approssimandolo con una retta:

$$\log(1 + x) \simeq a_0 + a_1 x. \quad (25.38)$$

Per $x = 0$ abbiamo che $a_0 = \log 1 = 0$. Perciò possiamo scrivere che

$$\log(1 + x) \simeq a_1 x, \quad (25.39)$$

che implica che

$$1 + x = e^{a_1 x}. \quad (25.40)$$

Ora, anche nel caso dell'esponenziale si può scrivere

$$e^x \simeq b_0 + b_1 x \quad (25.41)$$

e per $x = 0$ si trova che $\exp 0 = 1 = b_0$ e perciò

$$e^x \simeq 1 + b_1 x. \quad (25.42)$$

Dividiamo questa equazione per l'equazione (25.40) per ottenere

$$\frac{1 + b_1 x}{1 + x} = \frac{e^x}{e^{a_1 x}} = e^{(1-a_1)x}. \quad (25.43)$$

Per $a_1 = 1$ si ha $b_1 = 1$. Questo è l'unico caso in cui la relazione sopra scritta vale per ogni x e deve dunque essere generale. Possiamo quindi scrivere che

$$\log(1 + \alpha x) \simeq \alpha x \quad (25.44)$$

e

$$e^{\alpha x} \simeq 1 + \alpha x. \quad (25.45)$$

che, grazie all'equazione (25.44), possiamo scrivere come

La somma di piccoli incrementi relativi

In fisica è utile saper calcolare la somma di piccoli incrementi relativi di una grandezza fisica x . Se indichiamo che Δx l'incremento di una grandezza fisica x , il suo incremento relativo vale $\Delta x/x$. Ad esempio, se consideriamo un'auto che nel percorrere una strada a velocità costante consuma carburante, è abbastanza ragionevole pensare che la quantità di carburante bruciata per unità di tempo sia proporzionale alla massa della macchina M : più pesa la macchina più carburante si consuma. Ma la massa M dell'auto non resta costante, proprio perché una parte del carburante viene bruciata e dispersa. Se scriviamo $M = M_0 + m_i$ dove M_0 è la massa a vuoto dell'auto e m_i quella del carburante al tempo t_i possiamo dire che la massa di carburante $\Delta\mu$ consumato per unità di tempo Δt , $\Delta\mu/\Delta t$, al tempo t_i sarà proporzionale a m_i :

$$\frac{\Delta\mu}{\Delta t} \propto m_i, \quad (25.46)$$

cioè $\Delta\mu = \alpha m_i \Delta t$, con α costante. Ma $\Delta\mu$ è proprio di quanto diminuisce m_i , quindi possiamo scriverlo come $\Delta\mu = -\Delta m$, cioè come la variazione (che è negativa perché si tratta di una diminuzione) della massa del carburante presente. Sarà quindi

$$\frac{\Delta m}{m_i} = -\alpha \Delta t. \quad (25.47)$$

Il rapporto $\Delta m/m_i$ rappresenta la variazione relativa di carburante consumato per unità di tempo. Quanto vale la somma sulla sinistra? Per trovarlo osserviamo che si può scrivere sempre che

$$\log(m_i + \Delta m) = \log\left(m_i \left(1 + \frac{\Delta m}{m_i}\right)\right) \quad (25.48)$$

e dunque, sfruttando le proprietà dei logaritmi, che

$$\log(m_i + \Delta m) = \log m_i + \log\left(1 + \frac{\Delta m}{m_i}\right) \quad (25.49)$$

$$\log(m_i + \Delta m) \simeq \log m_i + \frac{\Delta m}{m_i}. \quad (25.50)$$

Quindi il rapporto $\Delta m_i/m$ si può riscrivere come

$$\frac{\Delta m}{m_i} \simeq \log(m_i + \Delta m) - \log m_i. \quad (25.51)$$

Ora sommiamo il primo e il secondo membro su più intervalli di tempo Δt , cioè calcoliamo

$$\sum_{i=1}^N \frac{\Delta m}{m_i} \simeq \sum_{i=1}^N \log(m_i + \Delta m) - \log m_i. \quad (25.52)$$

La somma S a destra si scrive come segue:

$$S = \log(m_0 + \Delta m) - \log m_0 + \log(m_1 + \Delta m) - \log m_1 + \dots \quad (25.53)$$

ma $m_{i+1} = m_i + \Delta m$ e quindi, in particolare, $m_1 = m_0 + \Delta m$; sostituendo

$$S = \log(m_0 + \Delta m) - \log m_0 + \log(m_0 + 2\Delta m) - \log m_0 + \dots \quad (25.54)$$

e in particolare i due logaritmi di $m_0 + \Delta m$ si cancellano e si ottiene

$$S = -\log m_0 + \log(m_0 + 2\Delta m) + \dots \quad (25.55)$$

Lo stesso accade a tutte le coppie di logaritmi intermedie dopo aver sommato N addendi resta

$$S = \log m_0 + N\Delta m - \log m_0 = \log m_N - \log m_0 = \log \frac{m_N}{m_0} \quad (25.56)$$

avendo indicato con m_N la massa m raggiunta dopo N passi. Tutte le volte che troviamo una somma di termini del tipo $\Delta x/x$ possiamo porla uguale al logaritmo del rapporto tra il valore che la grandezza

fisica assume al termine della somma x_f e quello x_i che aveva prima di iniziare la somma:

$$\sum_i \frac{\Delta x}{x_i} = \log \frac{x_f}{x_i}. \quad (25.57)$$

In analisi matematica la somma S diventa un **integrale** e si ha che

$$\int_x^b \frac{dx}{x} = \log \frac{b}{a}. \quad (25.58)$$

Funzioni trigonometriche

Scriviamo le funzioni seno e coseno approssimando ciascuna con un polinomio di primo grado:

$$\begin{aligned} \sin x &\simeq a_0 + a_1 x; \\ \cos x &\simeq b_0 + b_1 x. \end{aligned} \quad (25.59)$$

Imponendo che, per $x = 0$, $\sin x = 0$ e $\cos x = 1$ si trova che $a_0 = 0$ e $b_0 = 1$, dunque

$$\begin{aligned} \sin x &\simeq a_1 x; \\ \cos x &\simeq 1 + b_1 x. \end{aligned} \quad (25.60)$$

Ora imponiamo che $\sin^2 x + \cos^2 x = 1$:

$$1 = (a_1 x)^2 + (1 + b_1 x)^2 = (a_1^2 + b_1^2)x^2 + 1 + 2b_1 x \quad (25.61)$$

e trascurando il termine proporzionale a x^2 otteniamo che dev'essere $b_1 = 0$. Ora scriviamo il rapporto

$$R = \frac{\sin(x + \delta) - \sin x}{\delta}. \quad (25.62)$$

Quanto più δ diventa piccolo, tanto più il rapporto (detto **rapporto incrementale**) si avvicina al valore $R = 1$. Ora scriviamo le funzioni trigonometriche approssimandole con polinomi di primo grado:

$$R \simeq \frac{a_1(x + \delta) - a_1 x}{\delta} = a_1. \quad (25.63)$$

Ma poiché tale rapporto deve valere 1, almeno per δ molto piccoli, $a_1 = 1$. In definitiva

$$\begin{aligned} \sin x &\simeq x; \\ \cos x &\simeq 1; \\ \tan x &= \frac{\sin x}{\cos x} \simeq \sin x \simeq x; \end{aligned} \quad (25.64)$$

Altre funzioni

Consideriamo la funzione $f(x) = (r + x)^{-1}$, con r costante. Approssimando la funzione con una retta in prossimità di $x = 0$ si ottiene

$$\frac{1}{r + x} \simeq a_0 + a_1 x. \quad (25.65)$$

Per $x = 0$ abbiamo che $a_0 = 1/r$, quindi

$$\frac{1}{r + \varepsilon} \simeq \frac{1}{r} + a_1 \varepsilon, \quad (25.66)$$

dalla quale ricaviamo

$$a_1 = \frac{1}{\varepsilon} \left(\frac{1}{r + \varepsilon} - \frac{1}{r} \right) = \frac{1}{\varepsilon} \left(\frac{-\varepsilon}{r(r + \varepsilon)} \right) \simeq -\frac{1}{r^2}. \quad (25.67)$$

Quindi possiamo scrivere che, nell'intorno di $x = 0$,

$$\frac{1}{r + x} \simeq \frac{1}{r} - \frac{x}{r^2} \quad (25.68)$$

Equazioni differenziali a variabili separabili

Un'equazione differenziale è un'equazione nella quale compare una variabile e una sua variazione. Se, ad esempio, la grandezza fisica x varia da $x(0)$ a $x(T)$ in un tempo che va da $t = 0$ a $t = T$, la variazione $\Delta x = x(T) - x(0)$ che avviene nel tempo $\Delta t = T - 0 = T$, è legata al valore di x da un'equazione differenziale per cui

$$\Delta x = f(x(t), \Delta t) \quad (25.69)$$

dove f è una funzione sia di Δt che di x , che a sua volta dipende da t . Un'equazione differenziale a variabili separabili è un'equazione del tipo

$$\Delta x = \alpha x \Delta t. \quad (25.70)$$

Quest'equazione si dice *a variabili separabili* perché possiamo dividere entrambi i membri per x e ottenere

$$\frac{\Delta x}{x} = \alpha \Delta t \quad (25.71)$$

in cui a primo membro compaiono solo x e la sua variazione, mentre a secondo membro compare solo t . La soluzione di questa equazione è

$$x(t) = X_0 e^{\alpha t}. \quad (25.72)$$

dove X_0 è una costante. Infatti $x(0) = X_0$ e $x(T) = X_0 e^{\alpha T}$, per cui

$$\Delta x = X_0 (e^{\alpha T} - 1). \quad (25.73)$$

e quindi

$$\frac{\Delta x}{x} = \frac{X_0 (e^{\alpha T} - 1)}{X_0} = e^{\alpha T} - 1. \quad (25.74)$$

Se T è abbastanza piccolo, possiamo scrivere $e^{\alpha T} \simeq 1 + \alpha T$, infatti, per $T = 0$, l'esponenziale vale 1 e per $T > 0$, ma piccolo, è poco più grande. Sostituendo abbiamo

$$\frac{\Delta x}{x} \simeq 1 + \alpha T - 1 = \alpha T = \alpha \Delta t \quad (25.75)$$

che è proprio l'equazione che volevamo risolvere.

Unità naturali

La scelta delle **grandezze fisiche fondamentali** è del tutto arbitraria. Nel **Sistema Internazionale** (SI) le grandezze fondamentali sono quelle di **lunghezza**, **massa** e **tempo**. In questo sistema la **velocità** è una **grandezza fisica derivata**.

Non sempre questa è la scelta più opportuna. In effetti possiamo scegliere di considerare fondamentale la velocità e definire l'unità di misura come la velocità della luce c che, in questo sistema, evidentemente vale $c = 1$, che possiamo considerare adimensionale (le velocità infatti si misureranno in frazioni della velocità della luce). Se oltre alla velocità scegliamo l'energia E come altra grandezza fisica fondamentale, che possiamo misurare in un'unità a piacere, come il MeV o il GeV, tutte le altre si possono esprimere come grandezze fisiche derivate. Questo sistema di unità di misura si chiama delle **unità naturali**. In effetti, l'energia di una particella ferma è, secondo la teoria della relatività

$$E = mc^2 \quad (25.76)$$

ed essendo $c = 1$ e adimensionale, possiamo scrivere

$$E = m \quad (25.77)$$

dal che si vede che le masse si misurano in unità di energia. In effetti si dice che la massa del protone è di circa 1 GeV, intendendo dire che la massa in kg di questa particella vale circa $1 \text{ GeV} / (3 \times 10^8)^2 \text{ m}^2 \text{s}^{-2} = 1.6 \times 10^{-19} \times 10^9 \text{ J} \times (3 \times 10^8)^{-2} \text{ m}^{-2} \text{s}^2 \simeq 0.18 \times 10^{-26} \text{ Jm}^{-2} \text{s}^2$. Considerando che 1 J corrisponde a $1 \text{ kg m}^2 \text{s}^{-2}$ abbiamo che $m_p \simeq 1.8 \times 10^{-27} \text{ kg}$. Un elettrone, invece, *pesa* circa 0.5 MeV.

Se una particella è in moto la sua energia è tale che

$$E^2 = p^2 c^2 + m^2 c^4 \quad (25.78)$$

che, in unità naturali, si scrive

$$E^2 = p^2 + m^2. \quad (25.79)$$

È evidente da com'è scritta che quest'equazione, non solo la massa ha le stesse dimensioni fisiche dell'energia, ma anche le quantità di moto si misurano in queste unità.

Aggiungendo alle unità fondamentali quella definita da un'altra costante: la **costante di Planck** $\hbar = h/2\pi \simeq 6.6 \times 10^{-16} \text{ eVs}$ si ottiene il risultato secondo il quale le grandezze fisiche corrispon-

denti a lunghezza e tempo sono derivate e si misurano entrambe in unità di energia alla meno uno: $[L] = [E^{-1}]$, $[T] = [E^{-1}]$. Così un secondo diventa un tempo pari a

$$1 \text{ s} = \frac{1 \text{ s}}{\hbar \text{ eV s}} \simeq 0.1 \times 10^{16} \text{ eV}^{-1} = 0.1 \times 10^{10} \text{ MeV}^{-1} \quad (25.80)$$

e la lunghezza corrispondente a 1 m diventa

$$1 \text{ m} = \frac{1 \text{ m}}{3 \times 10^8 \text{ ms}^{-1} \times 6.6 \times 10^{-16} \text{ eV s}} \simeq 0.05 \times 10^8 \text{ eV}^{-1}. \quad (25.81)$$

Le unità naturali dunque si definiscono imponendo che $\hbar = c = 1$ come grandezze fisiche fondamentali e adimensionali. Sono molto comode per risolvere esercizi di cinematica relativistica, in cui tutte le potenze di c svaniscono, valendo 1 e non avendo dimensioni.

Prova a risolvere gli esercizi usando queste unità e vedrai che i conti si semplificheranno notevolmente. Puoi sempre tornare nelle unità canoniche moltiplicando per opportune potenze di c e $\hbar$ tali da riportare le unità di misura a quelle del SI.

Soluzione degli esercizi

Esercizio 11.1

osservati da Terra, gli orologi a bordo di un satellite GPS si muovono più lentamente perché la velocità della luce c è la stessa a bordo del satellite che si muove e per noi che siamo fermi. Espressa nel SI la velocità dei satelliti GPS è di 4 000 m/s. In unità di c è quindi

$$\beta = \frac{4 \times 10^3}{3 \times 10^8} \simeq 1.3 \times 10^{-5} \quad (\text{S.82})$$

e di conseguenza il fattore di Lorentz γ vale

$$\gamma = \frac{1}{\sqrt{1 - \beta^2}} \simeq 1.0000000001, \quad (\text{S.83})$$

differisce cioè da 1 per una parte su 10^{10} (possiamo scrivere $\gamma = 1 + 10^{-10}$). Di conseguenza, indicando con τ il tempo trascorso a bordo del satellite e con t quello misurato a Terra, essendo $t = \gamma\tau$, si ha che un secondo a bordo del satellite dura un po' più di quello a Terra: precisamente 10^{-10} secondi in più.

La durata di un giorno è maggiore di $86400 \times 10^{-10} \simeq 10^5 \times 10^{-10} = 10^{-5}$ secondi. In un mese gli orologi accumulano un ritardo pari a 30 volte questo, pari a 3×10^{-4} s e in un anno il ritardo ammonta a $12 \times 3 \times 10^{-4} = 0.0036$ s o 3.6 millesimi di secondo. Sembra poco, ma in questo tempo la luce percorre ben $3 \times 10^8 \times 3.6 \times 10^{-3} = 1\,080\,000$ m, cioè più di 1 000 km!

Dal momento che i satelliti GPS orbitano a una quota di 20 000 km, sbagliare la loro distanza di 1 000 km significa commettere un errore del 5 %, che naturalmente si riflette in un errore analogo sulla posizione rilevata a terra. Il che può significare un errore nella determinazione del punto d'intersezione delle sfere dell'ordine di alcuni km. Senza le correzioni relativistiche apportate agli orologi di

bordo (che a Terra sono starati in modo da marciare più rapidamente di quanto dovrebbero cosicché in volo sono sincronizzati con quelli a Terra), la costellazione GPS sarebbe del tutto inutile.

Esercizio 11.2

Il muone che viaggia verso la Terra vede quest'ultima venirgli incontro a una velocità $v = 0.995c$. Se fosse fermo la distanza tra lui e la Terra sarebbe di 10 km, ma dal momento che si vengono incontro, un raggio di luce che partisse dal muone e venisse riflesso dalla Terra impiegherebbe meno tempo a tornare al muone rispetto a quando il muone è fermo. Dal momento che le distanze si possono misurare valutando il tempo di andata e ritorno della luce, la costanza della sua velocità altera la misura delle distanze.

La distanza vista dal muone è $\ell' = \ell/\gamma$ dove γ è il fattore di Lorentz della Terra vista dal muone che vale

$$\gamma = \frac{1}{\sqrt{1 - \beta^2}} = \frac{1}{\sqrt{1 - (0.995)^2}} \simeq 10. \quad (\text{S.84})$$

La distanza da percorrere, secondo il muone, è dunque di appena 1 km.

Esercizio 11.3

Il fascio di N elettroni che attraversa la differenza di potenziale ΔV acquista un'energia cinetica complessiva $K = Ne\Delta V$, dove e è la carica dell'elettrone in modulo. Quando colpisce il bersaglio l'energia

cinetica del fascio è trasferita all'acqua. Il trasferimento di energia provoca un innalzamento della temperatura ΔT di m kg di questa per il quale

$$\Delta U = mc_a \Delta T \quad (\text{S.85})$$

essendo $c_a = 4.186 \text{ J/(kg K)}$ il calore specifico dell'acqua. Imponendo che $\Delta U = K$ si trova il numero di elettroni nel fascio:

$$N = \frac{mc_a \Delta T}{e \Delta V}. \quad (\text{S.86})$$

La massa dell'acqua si trova sapendo che la sua densità è $\rho = 1 \text{ kg}/\ell^3$ pertanto $m = 2 \text{ kg}$. Abbiamo dunque

$$N = \frac{2 \times 4.186 \times 10^{-4}}{1.6 \times 10^{-19} 10^5} \simeq 5 \times 10^{10}. \quad (\text{S.87})$$

L'energia totale di ogni elettrone è la somma della sua energia cinetica e della sua massa a riposo (che in unità naturali vale $511 \text{ keV} = 5.11 \times 10^5 \text{ eV}$). In queste unità, l'energia cinetica degli elettroni si scrive semplicemente come $K = 10^5 \text{ eV}$ e quindi

$$E = K + m \simeq 6 \times 10^5 \text{ eV}. \quad (\text{S.88})$$

Il fattore di Lorentz dell'elettrone è dunque

$$\gamma = \frac{E}{m} \simeq \frac{6}{5} \simeq 1.2 \quad (\text{S.89})$$

e di conseguenza la sua velocità β si ricava da

$$\beta^2 = 1 - \frac{1}{\gamma^2} \simeq 0.31 \quad (\text{S.90})$$

da cui si ricava che $v = c\beta \simeq 3 \times 10^8 \sqrt{0.31} \simeq 1.7 \times 10^8 \text{ m/s}$. Per percorrere i 150 m che separano il calorimetro dall'acceleratore dunque gli elettroni impiegano un tempo pari a

$$t = \frac{L}{v} = \frac{150}{1.7 \times 10^8} \simeq 9 \times 10^{-7} \text{ s} = 0.9 \mu\text{s}. \quad (\text{S.91})$$

Se non avessimo usato la cinematica relativistica, la velocità degli elettroni sarebbe risultata essere pari a

$$v = \sqrt{\frac{2K}{m}} \simeq \sqrt{\frac{2}{5}} \simeq 0.6, \quad (\text{S.92})$$

che in unità SI vale $v = 0.6 \times 3 \times 10^8 = 1.8 \times 10^8 \text{ m/s}$. Gli elettroni dovrebbero quindi impiegare

$$t_c = \frac{L}{v} \simeq \frac{150}{1.8 \times 10^8} = 0.8 \mu\text{s}. \quad (\text{S.93})$$

Eseguendo l'esperimento si trova, in effetti, che quando l'acqua aumenta la propria temperatura di $10^{-4} \text{ }^\circ\text{C}$, il tempo impiegato dagli elettroni a raggiungere il calorimetro è compatibile con l'essere $0.9 \mu\text{s}$ e non compatibile con $0.8 \mu\text{s}$.

Esercizio 11.4

I protoni circolano lungo un anello grazie alla [Forza di Lorentz](#) che si scrive

$$\mathbf{F} = e\mathbf{v} \wedge \mathbf{B}. \quad (\text{S.94})$$

Dal momento che il campo magnetico è perpendicolare a $\mathbf{v}$ il modulo della forza vale semplicemente $F = evB$. Nel sistema di riferimento del protone questa forza dev'essere assente e quindi dev'essere uguale e contraria alla forza centrifuga il cui modulo è $F = mv^2/r$. Uguagliando queste due quantità si ricava

$$evB = m \frac{v^2}{r} \quad (\text{S.95})$$

da cui possiamo ricavare il raggio dell'orbita

$$r = \frac{mv}{eB}. \quad (\text{S.96})$$

L'energia cinetica dei protoni è molto più alta della loro massa a riposo perciò possiamo porre $v \simeq c$. In effetti il fattore di Lorentz dei protoni, sapendo che $m \simeq 1 \text{ GeV}$, vale

$$\gamma = \frac{E}{m} \simeq \frac{7 \times 10^{12}}{10^9} = 7000, \quad (\text{S.97})$$

e quindi

$$\beta = \sqrt{1 - \frac{1}{\gamma^2}} \simeq 0.9999999898. \quad (\text{S.98})$$

Se non tenessimo conto della relatività otterremmo per r il valore di

$$r = \frac{1.7 \times 10^{-27} \times 3 \times 10^8}{1.6 \times 10^{-19} \times 8.4} \simeq 0.38 \text{ m} \quad (\text{S.99})$$

(la massa del protone in unità SI si trova dividendo la massa in GeV per c^2). In effetti la massa del protone aumenta di un fattore γ e dunque il raggio di curvatura dell'orbita è di

$$\gamma r = 7000 \times 0.38 = 2660 \text{ m}. \quad (\text{S.100})$$

La lunghezza della traiettoria è quindi di ben $L = 2\pi r \simeq 17 \text{ km}$! In effetti LHC è ancora più lungo: 27 km. Il motivo è che la macchina non ha la forma di una circonferenza, ma di un ottagono con gli spigoli curvi, per semplificarne la costruzione. Il raggio di curvatura che abbiamo calcolato, dunque, è quello che si trova in corrispondenza degli spigoli dell'ottagono.

Esercizio 13.1

Il tempo in prossimità di un corpo massivo scorre in maniera diversa rispetto a quanto fa in assenza di gravità. I satelliti della costellazione GPS orbitano a una quota di 20 000 km da Terra (la cui massa è di circa $6 \times 10^{24} \text{ kg}$) e sono quindi soggetti a un potenziale gravitazionale

$$\mathcal{G} = G \frac{M}{r} = 6.6 \times 10^{-11} \frac{6 \times 10^{24}}{2 \times 10^4} \simeq 2.0 \times 10^{10} \text{ m}^2 \text{ s}^{-2}. \quad (\text{S.101})$$

Un orologio sulla Terra invece è soggetto a un potenziale di

$$\mathcal{G}' = G \frac{M}{r'} = 6.6 \times 10^{-11} \frac{6 \times 10^{24}}{6 \times 10^3} \simeq 6.6 \times 10^{10} \text{ m}^2 \text{ s}^{-2}. \quad (\text{S.102})$$

Rispetto a un orologio molto lontano da ogni sorgente di campo gravitazionale il tempo a Terra scorre più lentamente: un secondo trascorso secondo il primo, per l'orologio a Terra dura

$$\frac{\mathcal{G}'}{c^2} \simeq 73 \mu\text{s} \quad (\text{S.103})$$

in più. A bordo dei satelliti invece il ritardo è di

$$\frac{\mathcal{G}}{c^2} \simeq 22 \mu\text{s} \quad (\text{S.104})$$

quindi gli orologi di bordo anticipano mediamente di $73 - 22 = 51 \mu\text{s}$ al secondo rispetto a quelli a Terra. Questa correzione non è costante (perché l'orbita non è circolare) e il suo valore è calcolato a Terra e inviato a bordo del satellite durante le comunicazioni con le stazioni di controllo. Notate che questa correzione compensa parzialmente quella dovuta alla relatività ristretta.

Esercizio 14.1

Data la lunghezza d'onda della luce, possiamo ricavarne la frequenza secondo la formula

$$\nu = \frac{c}{\lambda} \quad (\text{S.105})$$

facile da ricordare perché ν ha le dimensioni di un tempo alla meno uno e λ quelle di una lunghezza. Per ottenere una grandezza fisica con le dimensioni di $[T^{-1}]$ è dunque necessario dividere una velocità per una lunghezza. Ricaviamo perciò che la frequenza della luce del Sole è compresa tra $\nu_2 = 3 \times 10^8 / 400 \times 10^{-9} = 0.75 \times 10^{15} \text{ Hz}$ e $\nu_1 = 3 \times 10^8 / 750 \times 10^{-9} = 0.4 \times 10^{15}$. Sapendo che $h = 6.63 \times 10^{-34} \text{ Js}$ o $4.14 \times 10^{-15} \text{ eVs}$, l'energia dei fotoni corrispondenti è quindi compresa tra

$$E_1 = h\nu_1 = 4.14 \times 10^{-15} \times 0.4 \times 10^{15} = 1.66 \text{ eV} \quad (\text{S.106})$$

e

$$E_2 = h\nu_2 = 4.14 \times 10^{-15} \times 0.75 \times 10^{15} = 3.11 \text{ eV}. \quad (\text{S.107})$$

L'energia media trasportata dai fotoni è dunque attorno ai 2 eV. L'energia che hanno gli elettroni che

hanno subito effetto fotoelettrico è pari a quella dei fotoni sottratta dell'energia di legame dell'elettrone. Nel caso del silicio, l'energia necessaria a un elettrone per diventare conduttore è di $1.1 - 1.2$ eV. Il silicio è il materiale usato per costruire le celle fotovoltaiche. La radiazione solare che incide su una cella libera gli elettroni del silicio e li fa passare in uno stato in cui sono praticamente liberi. Il moto di questi elettroni genera la corrente fotovoltaica.

Esercizio 14.2

Un elettrone accelerato da una differenza di potenziale di 80 kV ha un'energia cinetica pari a 80 keV: il 16 % della sua energia a riposo di 511 keV. La sua energia totale è dunque $80 + 511 = 591$ keV (è la somma dell'energia cinetica e di quella a riposo). Possiamo quindi ricavarne la velocità sfruttando la definizione del fattore di Lorentz

$$\gamma^2 = \frac{1}{1 - \beta^2} \quad (\text{S.108})$$

da cui, sapendo che $\gamma = E/m \simeq 1.16$ si ricava che

$$\beta^2 = 1 - \frac{1}{\gamma^2} \simeq 1 - \frac{1}{1.16^2} \simeq 0.26 \quad (\text{S.109})$$

e quindi $\beta \simeq 0.51$ (in altre parole gli elettroni del microscopio viaggiano a circa metà della velocità della luce. La quantità di moto degli elettroni è dunque pari a $p = \gamma m \beta = 1.16 \times 511 \times 0.51 \simeq 300$ keV e di conseguenza sono assimilabili a onde di lunghezza d'onda pari a

$$\lambda = \frac{h}{p} = \frac{2\pi}{3 \times 10^5} \simeq 2.1 \times 10^{-5} \text{ eV}^{-1} \quad (\text{S.110})$$

(ricordiamo che in unità naturali $\hbar = h/2\pi = 1$ e le distanze si misurano in unità di energia alla meno uno). Per avere la lunghezza in metri basta moltiplicare per opportune potenze di c e di $\hbar$ in unità SI. Moltiplicando per $\hbar$ che ha le dimensioni di un'energia per un tempo otteniamo una grandezza che le

dimensioni di un tempo, che possiamo trasformare in una lunghezza moltiplicando per una velocità:

$$\begin{aligned} \lambda[\text{m}] &= \lambda[\text{eV}^{-1}] \hbar[\text{eV s}] c[\text{ms}^{-1}] \\ &= 2.1 \times 10^{-5} \times 6.58 \times 10^{-16} \times 3 \times 10^8 \simeq 4.3 \times 10^{-11} \end{aligned} \quad (\text{S.111})$$

Per confronto le onde luminose hanno una lunghezza d'onda dell'ordine dei 500 nm, cioè di 5×10^{-7} m, vale a dire 5 ordini di grandezza maggiore! Un microscopio elettronico perciò permette di apprezzare dettagli 100 000 volte più piccoli rispetto a quello ottico.

Esercizio 15.1

Se il guadagno di energia ΔE da parte di una particella di energia E che attraversa in un tempo Δt una regione di spazio in cui viene accelerata è

$$\Delta E = \alpha E \Delta t \quad (\text{S.112})$$

dividendo per E abbiamo che

$$\frac{\Delta E}{E} = \alpha \Delta t. \quad (\text{S.113})$$

Quando si trova un'equazione del genere, la [soluzione](#) è sempre del tipo

$$E = E_0 e^{\alpha t}. \quad (\text{S.114})$$

dove E_0 è l'energia della particella al tempo $t = 0$ ed E quella al tempo t . Poiché α deve avere le dimensioni fisiche di un tempo alla meno uno, dovendo essere l'argomento dell'esponenziale adimensionale, possiamo definire $\alpha = 1/\tau$ e scrivere

$$E = E_0 e^{\frac{t}{\tau}}, \quad (\text{S.115})$$

dove τ è una grandezza che ha le dimensioni fisiche di un tempo $[\tau] = [T]$.

Per calcolare la probabilità di aver viaggiato per un tempo t senza perdere energia, facciamo così: supponiamo che una particella di energia E_0 viaggi per un tempo t_0 sufficientemente piccolo durante il quale la probabilità di guadagnare energia si

può considerare costante: $P \simeq p$. Se dopo essere sopravvissuta all'attraversamento di questa regione, avendo guadagnato energia, si muove ancora per un tempo t , la probabilità di guadagnare energia dall'istante $t = 0$ è ora

$$P = p^2 \quad (\text{S.116})$$

e così via: dopo aver attraversato n strati, ciascuno per una durata molto breve, la probabilità di sopravvivenza è $P = p^n$. Scriviamo n come $n = t/t_0$ e osserviamo che P deve essere un numero compreso tra 0 e 1, così come p . Ridefinendo $p = 1/k$, possiamo scrivere

$$P = k^{-\frac{t}{t_0}} \quad (\text{S.117})$$

Così facendo è evidente che, per $t = 0$, $P = 1$ (cioè la probabilità di sopravvivenza dopo un tempo $t = 0$ è 1, che è ovvio perché non essendo trascorso alcun tempo lo stato della particella non può essere cambiato), mentre per $t \rightarrow \infty$, $P \rightarrow 0$ (cioè per tempi molto molto lunghi è improbabile che la particella sopravviva). Questo sembra del tutto ragionevole.

Possiamo naturalmente scrivere che

$$P = k^{-\frac{t}{t_0} \frac{\tau}{\tau}} = \left(k^{\frac{t}{\tau}}\right)^{-\frac{\tau}{t_0}} = \left(k^{\frac{t}{\tau}}\right)^{-\frac{\tau}{t_0}} E_0^{\frac{\tau}{t_0}} E_0^{-\frac{\tau}{t_0}} \quad (\text{S.118})$$

Portando dentro la parentesi $E_0^{-\frac{\tau}{t_0}}$ e sostituendo la stessa con l'equazione (S.115) si ottiene

$$P = \left(E_0 e^{\frac{t}{\tau}}\right)^{-\frac{\tau}{t_0}} E_0^{\frac{\tau}{t_0}} = E^{-\frac{\tau}{t_0}} E_0^{\frac{\tau}{t_0}}. \quad (\text{S.119})$$

Il fattore $E_0^{\frac{\tau}{t_0}}$ è una costante

$$E_0^{\frac{\tau}{t_0}} = A \quad (\text{S.120})$$

così come l'esponente di E

$$\frac{\tau}{t_0} = \gamma \quad (\text{S.121})$$

perciò

$$P = A E^{-\gamma}. \quad (\text{S.122})$$

Quindi il numero di particelle di energia E che giungeranno a noi dopo aver viaggiato per un tempo t è proporzionale a questa probabilità e va dunque come una legge di potenza.

Esercizio 16.1

L'energia di una particella di massa m e quantità di moto p è

$$E = \sqrt{p^2 c^2 + m^2 c^4} \quad (\text{S.123})$$

dove c è la velocità della luce.

Considerando che nello stato iniziale abbiamo un neutrone fermo, indicando con m_n la sua massa, l'energia vale

$$E_{in} = m_n c^2. \quad (\text{S.124})$$

Nello stato finale, per la conservazione della quantità di moto, deve essere

$$\vec{p}_p = -\vec{p}_e \quad (\text{S.125})$$

dove $\vec{p}_p$ indica la quantità di moto del protone e $\vec{p}_e$ quella dell'elettrone. Abbiamo che

$$|\vec{p}_p| = |\vec{p}_e| = p \quad (\text{S.126})$$

e quindi possiamo scrivere che

$$E_{fin} = \sqrt{p^2 c^2 + m_p^2 c^4} + \sqrt{p^2 c^2 + m_e^2 c^4}. \quad (\text{S.127})$$

Imponendo che l'energia iniziale sia uguale a quella finale possiamo scrivere che

$$m_n^2 c^2 = 2p^2 + m_p^2 c^2 + m_e^2 c^2 + 2\sqrt{p^2 + m_p^2 c^2} \sqrt{p^2 + m_e^2 c^2}. \quad (\text{S.128})$$

Lasciamo a secondo membro solo il prodotto delle radici:

$$(m_n^2 - m_p^2 - m_e^2) c^2 - 2p^2 = 2\sqrt{p^2 + m_p^2 c^2} \sqrt{p^2 + m_e^2 c^2} \quad (\text{S.129})$$

ed eleviamo al quadrato entrambi i membri:

$$(M^2 c^2 - 2p^2)^2 = 4(p^2 + m_p^2 c^2)(p^2 + m_e^2 c^2), \quad (\text{S.130})$$

dove $M^2 = (m_n^2 - m_p^2 - m_e^2)$. Notate che l'espressione a secondo membro ha le dimensioni di una massa al quadrato e così, sebbene si potesse definire quest'espressione con una generica variabile x , abbiamo preferito usare il simbolo M^2 per rendere palese il fatto che si tratta di una grandezza fisica che ha queste dimensioni. Questo aiuta sempre nella soluzione di un problema di fisica.

Sviluppiamo i quadrati e semplifichiamo. Prima abbiamo

$$M^4 c^4 + 4p^4 - 4p^2 M^2 c^2 = 4p^4 + 4p^2 (m_e^2 c^2 + m_p^2 c^2) + 4(m_p m_e)^2 c^4 \quad (\text{S.131})$$

e quindi, semplificando il termine $4p^4$ presente a destra e a sinistra, portando i termini proporzionali a p^2 a primo membro e gli altri a secondo membro e, infine, dividendo tutto per c^2 otteniamo

$$4p^2 (m_e^2 + m_p^2 + M^2) = M^4 c^2 - 4(m_p m_e)^2 c^2. \quad (\text{S.132})$$

Si vede subito che, avendo usato il simbolo M^2 per rappresentare la combinazione presente nell'equazione (S.129), l'espressione trovata è dimensionalmente corretta. Infatti, a primo membro abbiamo una quantità di moto al quadrato, che ha le dimensioni di una massa al quadrato per una velocità al quadrato, moltiplicata per una cosa che ha le dimensioni di una massa al quadrato a sua volta. Pertanto il primo membro ha le dimensioni di una massa alla quarta per una velocità al quadrato. Che sono le dimensioni fisiche del secondo membro, come appare in maniera evidente dalla forma dell'equazione.

A questo punto è facile ricavare p^2 come

$$p^2 = \frac{M^4 - 4(m_p m_e)^2}{4(M^2 + m_e^2 + m_p^2)} c^2 = \frac{M^4 - 4(m_p m_e)^2}{4m_n^2} c^2. \quad (\text{S.133})$$

Facciamo un ultimo controllo dimensionale: il numeratore della frazione ha le dimensioni di una massa alla quarta e il denominatore di una massa al quadrato. Quindi la frazione ha le dimensioni di una massa al quadrato che, moltiplicata per una velocità al quadrato (c^2) dà proprio le dimensioni di una quantità di moto al quadrato.

Si vede subito che la quantità di moto p è una costante che dipende solo dalle masse delle particelle partecipanti alla reazione: m_n , m_p e m_e . Trascurando m_e rispetto alle altre due masse (questa è un'approssimazione ragionevole dal momento che $m_e \simeq m_p/2000$) e sapendo che $m_n \simeq m_p$ possiamo riscrivere questa quantità di moto come

$$p^2 \simeq \frac{(m_n^2 - m_p^2)^2}{4m_n^2} c^2 = \frac{(m_n - m_p)^2 (m_n + m_p)^2}{4m_n^2} c^2 \simeq \frac{(m_n - m_p)^2 (2m_n)^2}{4m_n^2} c^2 = (m_n - m_p)^2 c^2. \quad (\text{S.134})$$

La massima quantità di moto permessa per un elettrone da decadimento β è dunque dell'ordine di $p_{max} \simeq (m_n - m_p)c$. L'energia di un elettrone con questa quantità di moto è

$$E_{max} = \sqrt{(m_n - m_p)^2 c^4 + m_e^2 c^4} = c^2 \sqrt{(m_n - m_p)^2 + m_e^2} \quad (\text{S.135})$$

Sostituendo i valori $m_n \simeq 1.675 \times 10^{-27}$ kg, $m_p \simeq 1.673 \times 10^{-27}$ kg, $m_e \simeq 9 \times 10^{-31}$ e $c = 3 \times 10^8$ m/s abbiamo che

$$E_{max} \simeq 2 \times 10^{-13} \text{ J} \quad (\text{S.136})$$

(l'unità di misura è quella del SI, avendo usato tutte unità di questo sistema). Possiamo riscrivere questo valore in un'unità più adeguata per la fisica delle particelle: in eV. Per farlo basta ricordare che $1 \text{ eV} = 1.6 \times 10^{-19} \text{ J}$ e quindi, trattando le unità di misura come se fossero quantità algebriche, abbiamo che

$$1 = \frac{1 \text{ eV}}{1.6 \times 10^{-19} \text{ J}} \quad (\text{S.137})$$

così che, moltiplicando per 1 l'espressione in (S.136) otteniamo

$$E_{max} \simeq 2 \times 10^{-13} \text{ J} \times \frac{1 \text{ eV}}{1.6 \times 10^{-19} \text{ J}} \simeq 10^6 \text{ eV} = 1 \text{ MeV} . \quad (\gamma mv)^2 \simeq \left(\frac{63 \times 10^6 \times 1.6 \times 10^{-19}}{3 \times 10^8} \right)^2 - (9 \times 10^{-31} \cdot 3 \times 10^8)^2 . \quad (\text{S.144})$$

Sperimentalmente è proprio quello che si osserva, il che significa che la massa del neutrino è trascurabile. Il valore massimo dell'energia dell'elettrone varia, anche se di poco, da specie atomica a specie atomica a causa della diversa energia di legame con la quale i neutroni sono legati nel nucleo.

Esercizio 16.2

Una particella di massa m e carica q che si muove con velocità v perpendicolarmente alle linee di forza di un campo magnetico d'intensità B percorre una traiettoria circolare di raggio

$$r = \gamma \frac{mv}{qB} , \quad (\text{S.139})$$

dove

$$\gamma = \frac{1}{\sqrt{1 - \left(\frac{v}{c}\right)^2}} . \quad (\text{S.140})$$

Possiamo ricavare B invertendo la formula del raggio:

$$B = \gamma \frac{mv}{qr} . \quad (\text{S.141})$$

Per ricavarne il valore dobbiamo conoscere, oltre a $m \simeq 9 \times 10^{-31} \text{ kg}$ (la massa del positrone, identica a quella dell'elettrone) e $q = 1.6 \times 10^{-19} \text{ C}$ (l'opposto della carica elettrica dell'elettrone), il valore di v .

Questo possiamo ottenerlo sapendo che

$$E^2 = p^2 c^2 + m^2 c^4 = (\gamma mv)^2 c^2 + m^2 c^4 \quad (\text{S.142})$$

per cui

$$\gamma mv = \sqrt{\frac{E^2}{c^2} - m^2 c^2} . \quad (\text{S.143})$$

Sappiamo che, in un caso $E = 63 \text{ MeV}$ pertanto avremo, in unità del [SI](#),

Si vede subito che il secondo addendo sotto la radice è trascurabile rispetto al primo essendo l'ordine di grandezza del primo addendo pari a $(6 - 19 - 8) \times 2 = -42$, mentre per il secondo abbiamo $(-31 + 8) \times 2 = -46$. Il risultato è che possiamo estrarre la radice solo del primo addendo per cui

$$\gamma mv \simeq 33.6 \times 10^{-21} \quad (\text{S.145})$$

in unità del [SI](#) (kg m s^{-1}). Quando $E = 23 \text{ MeV}$, evidentemente la quantità di moto vale, nelle stesse unità, 12.4×10^{-21} .

Per valutare B dobbiamo conoscere i rispettivi raggi di curvatura, che possiamo misurare sull'immagine con un righello, con una risoluzione di 1 mm, sapendo che lo spessore dell'assorbitore è di 6 mm.

Possiamo misurare la corda più lunga possibile della traccia meno energetica (quella più curva), che è $L = 43 \pm 1 \text{ mm}$. La **sagitta**, cioè la distanza tra la corda e il punto medio dell'arco di circonferenza corrispondente, è $s = 5 \pm 1 \text{ mm}$. Utilizzando il teorema di Talete possiamo scrivere che

$$2 \frac{r - s}{L} = \frac{L}{2s} \quad (\text{S.146})$$

da cui possiamo ricavare r come

$$r = \frac{L^2}{4s} + s \simeq 97 \text{ mm} . \quad (\text{S.147})$$

L'errore da attribuire a questa quantità si dovrebbe ricavare usando le regole della [propagazione degli errori](#). In questo caso possiamo limitarci a stimarlo osservando che l'errore relativo sulla misura della sagitta (il peggiore tra i due) è dell'ordine del 20 % ($\delta s/s = 1/5 = 0.2$). Il valore di r non può essere noto con una precisione migliore di questa quindi $\delta r \simeq 19 \text{ mm}$ e $r = (97 \pm 19) \text{ mm}$. Da questa misura ricaviamo il valore del campo che è

$$B \simeq \frac{2.4 \times 10^{-21}}{1.6 \times 10^{-19} \cdot 97 \times 10^{-3}} \simeq 0.15 \text{ T} . \quad (\text{S.148})$$

L'unità di misura di B è ovviamente quella del sistema internazionale, avendo usato tutte unità in questo sistema. Anche in questo caso, l'errore relativo $\delta B/B$ deve essere al meglio del 20 %, quindi deve valere $\delta B \simeq 0.03$ T: $B = (0.15 \pm 0.03)$ T. Questo valore coincide, entro gli errori, con quello pubblicato nell'articolo di Anderson [22] di 15 000 Gauss.

Poiché l'energia iniziale del positrone è un fattore circa 3 rispetto a quella finale, il raggio di curvatura prima di attraversare l'assorbitore deve essere dello stesso fattore più ampio. Ricaviamo la sagitta in funzione del raggio e della corda:

$$s = \frac{r \pm \sqrt{(r-L)(r+L)}}{2} \quad (\text{S.149})$$

(ci sono due possibili valori della sagitta; quello che interessa a noi è il più piccolo). Nel caso in cui $L \ll r$ possiamo scrivere

$$r^2 - L^2 = r^2 \left(1 - \left(\frac{L}{r}\right)^2\right) \quad (\text{S.150})$$

e quindi, essendo $\sqrt{1-x} \simeq 1 - \frac{x}{2}$,

$$\sqrt{(r-L)(r+L)} \simeq r \sqrt{1 - \left(\frac{L}{r}\right)^2} \simeq r \left(1 - \frac{1}{2} \left(\frac{L}{r}\right)^2\right) \quad (\text{S.151})$$

e, in definitiva

$$s \simeq \frac{L^2}{2r}. \quad (\text{S.152})$$

Quando r diventa un fattore tre più ampio, s diventa un fattore 3 più piccola, quindi ci aspettiamo una sagitta di circa 1.7 mm, che è proprio quel che si vede eseguendo la misura (con la risoluzione che abbiamo assunto possiamo vedere che la sagitta è di circa 2 mm). Evidentemente il positrone proviene dal lato in cui la sua energia è maggiore, quindi, guardando la foto si comprende che percorre gli archi di circonferenza in senso antiorario.

La **Forza di Lorentz** $\vec{F} = q\vec{v} \wedge \vec{B}$ è diretta in modo tale da essere perpendicolare alla velocità e al campo e il suo verso deve essere quello uscente dal

palmo della mano destra con il pollice messo in direzione della velocità e le altre dita in direzione del campo. Il campo quindi deve essere perpendicolare al piano della foto ed entrante in esso.

Esercizio 17.1

La prima reazione $p + p \rightarrow p + \bar{p}$ non è permessa perché non conserva il numero barionico (che vale 2 nello stato iniziale e 0 in quello finale). La reazione $p + p \rightarrow p + p + p + \bar{p}$ invece è consentita per il motivo opposto.

La produzione di un neutrone e un positrone dallo scontro tra un protone e un pione neutro è vietata perché non è conservato il numero leptonico (che nello stato iniziale è nullo e nello stato finale vale -1).

La reazione $p + n \rightarrow n + p + \pi^0$ può avvenire perché non viola alcuna legge di conservazione.

Il decadimento del pione carico in $\pi^0 + \nu_e$ non può avvenire perché non conserva il numero leptonico.

La reazione $\pi^- + p \rightarrow n + \pi^0$ è consentita perché tutti i numeri quantici sono conservati.

Lo scontro tra un muone e un neutrone non può dare origine a un neutrone un protone e un anti-neutrino perché se il numero leptonico è conservato ($L_\mu(\mu^+) = L_\mu(\bar{\nu}_\mu)$), non lo è quello barionico (c'è un solo barione, il neutrone, con $B = 1$ a sinistra, e due barioni con $B = 1$, a destra).

Per motivi analoghi è vietata la reazione $\mu^- + e^- \rightarrow \bar{p} + e^-$. A destra abbiamo un barione, che non c'è a sinistra. Inoltre abbiamo, nello stato iniziale, $L_e = 1$ e $L_\mu = 1$, mentre nello stato finale $L_e = 1$, ma $L_\mu = 0$.

Lo scontro tra due pioni può dare origine a una coppia protone-antiprotone, perché il numero barionico è complessivamente nullo nello stato finale, così come nello stato iniziale.

Non può invece dare origine alla produzione di $p + n + \pi^-$ perché il numero barionico dello stato finale vale $B = 2$, mentre è sempre nullo nello stato iniziale.

Esercizio 18.1

Consideriamo l'urto di un protone di massa m e con velocità v con un altro protone fermo nel sistema di riferimento del laboratorio e che nell'urto si produca un π^0 . La reazione che stiamo studiando è

$$p + p \rightarrow p + p + \pi^0. \quad (\text{S.153})$$

Notate che tutti i numeri quantici sono conservati (carica elettrica, [numero barionico](#), [numero leptonico](#)). La quantità di moto iniziale è

$$P_{in} = \gamma m v. \quad (\text{S.154})$$

dove γ è il [fattore di Lorentz](#) che fa *aumentare* la massa efficace della particella, man mano che la sua velocità si avvicina a quella della luce. Nello stato finale la quantità di moto è

$$P_{out} = \gamma_1 m v_1 + \gamma_2 m v_2 + \gamma_\pi m_\pi v_\pi \quad (\text{S.155})$$

dove le grandezze con il pedice 1 si riferiscono a un protone e quelle col pedice 2 all'altro. Oltre alla quantità di moto deve essere conservata l'energia, perciò abbiamo che

$$\begin{aligned} \sqrt{\gamma^2 m^2 v^2 c^2 + m^2 c^4} + m c^2 &= \sqrt{\gamma_1^2 m^2 v_1^2 c^2 + m^2 c^4} + \\ &\quad \sqrt{\gamma_2^2 m^2 v_2^2 c^2 + m^2 c^4} + \\ &\quad \sqrt{\gamma_\pi^2 m_\pi^2 v_\pi^2 c^2 + m_\pi^2 c^4}. \end{aligned} \quad (\text{S.156})$$

Usando questo approccio vediamo subito che i conti si complicano non poco. In questi casi conviene usare un *trucco*. Innanzi tutto osserviamo che

$$E^2 - p^2 c^2 = m^2 c^4. \quad (\text{S.157})$$

Questa relazione vale per qualsiasi particella di massa m , quantità di moto p ed energia E . La massa a riposo di una particella è una costante, che non dipende dal sistema di riferimento nel quale si misura. Questo significa che la combinazione $E^2 - p^2 c^2$ è a sua volta una costante indipendente dal sistema di riferimento.

Possiamo allora calcolare questa quantità nel sistema di riferimento che ci fa più comodo. Nel caso della reazione in esame l'energia iniziale, misurata nel sistema del laboratorio in cui uno dei protoni è fermo e l'altro ha quantità di moto p , vale

$$E_{in} = \sqrt{p^2 c^2 + m^2 c^4} + m c^2 \quad (\text{S.158})$$

mentre la quantità di moto totale vale semplicemente

$$P_{in} = p. \quad (\text{S.159})$$

La differenza tra l'energia al quadrato e la quantità di moto al quadrato moltiplicata per c^2 è dunque

$$E_{in}^2 - P_{in}^2 c^2 = 2m^2 c^4 + 2m c^2 \sqrt{p^2 c^2 + m^2 c^4}. \quad (\text{S.160})$$

Per calcolare la stessa quantità nello stato finale, poiché possiamo scegliere il sistema di riferimento che ci fa più comodo, possiamo metterci nel sistema di riferimento del centro di massa, cioè in quello in cui le tre particelle risultano avere quantità di moto tali da annullarsi a vicenda, per cui

$$P_{fin} = 0. \quad (\text{S.161})$$

Nel caso più semplice le tre particelle sono prodotte ferme nel centro di massa e l'energia è semplicemente la somma delle masse moltiplicate per la velocità della luce al quadrato:

$$E_{fin} = 2m c^2 + m_\pi c^2. \quad (\text{S.162})$$

Attenzione! Il fatto che $P_{in} \neq P_{fin}$ non significa che stiamo violando la conservazione della quantità di moto! Le due quantità infatti sono calcolate in sistemi di riferimento diversi. Quello che importa è che la differenza $E_{fin}^2 - P_{fin}^2 c^2$ rimanga costante. Questa differenza vale

$$4m^2 c^4 + m_\pi^2 c^4 + 4m m_\pi c^4. \quad (\text{S.163})$$

Imponendo che questa sia uguale a quella calcolata nello stato iniziale (anche se in un sistema di riferimento diverso) otteniamo l'equazione

$$2m^2c^4 + m_\pi^2c^4 + 4mm_\pi c^4 = 2mc^2\sqrt{p^2c^2 + m^2c^4} \quad (\text{S.164})$$

da cui si ricava

$$\sqrt{p^2c^2 + m^2c^4} = E = mc^2 + \frac{m_\pi^2}{2m}c^2 + 2m_\pi c^2. \quad (\text{S.165})$$

Quella che abbiamo appena trovato è la minima energia (**energia di soglia**) che deve avere un protone affinché, scontrandosi con un protone fermo, possa dare luogo alla produzione di un π^0 .

Esercizio 18.2

Abbiamo due fotoni la cui quantità di moto è $\vec{p}_1$ e $\vec{p}_2$, con, rispettivamente, energie E_1 ed E_2 . Poiché i fotoni hanno massa nulla deve essere

$$E_1^2 - p_1^2c^2 = 0 \quad (\text{S.166})$$

e

$$E_2^2 - p_2^2c^2 = 0. \quad (\text{S.167})$$

Supponiamo che i fotoni siano il prodotto del decadimento di una particella di massa M . Nel sistema di riferimento in cui questa particella è ferma l'energia è Mc^2 e la quantità di moto è nulla. La differenza tra l'energia al quadrato e la quantità di moto al quadrato moltiplicata per c^2 , che è un invariante relativistico, vale dunque Mc^2 . Questa grandezza è appunto un invariante, perciò non cambia passando da un sistema di riferimento a un altro. Calcoliamola nel sistema di riferimento del laboratorio, in cui la particella di massa M si sta muovendo e decade nei due fotoni. Per la conservazione della quantità di moto abbiamo che

$$\vec{P}_M = \vec{p}_1 + \vec{p}_2 \quad (\text{S.168})$$

e per la conservazione dell'energia si deve avere che

$$E_M = E_1 + E_2. \quad (\text{S.169})$$

L'energia totale al quadrato nello stato finale è

$$E_M^2 = (E_1 + E_2)^2 = E_1^2 + E_2^2 + 2E_1E_2 \quad (\text{S.170})$$

mentre il prodotto $P_M^2c^2$ vale

$$P^2c^2 = (\vec{p}_1 + \vec{p}_2)^2c^2 = p_1^2c^2 + p_2^2c^2 + 2p_1p_2c^2\cos\theta \quad (\text{S.171})$$

dove θ è l'angolo formato dai vettori $\vec{p}_1$ e $\vec{p}_2$. La differenza $E_M^2 - P^2c^2$ è un invariante perciò deve valere Mc^2 :

$$E_1^2 + E_2^2 + 2E_1E_2 - (p_1^2c^2 + p_2^2c^2 + 2p_1p_2c^2\cos\theta) = Mc^2. \quad (\text{S.172})$$

Ora osserviamo che $E_1^2 - p_1^2c^2 = E_2^2 - p_2^2c^2 = 0$ e quindi

$$2E_1E_2 - 2p_1p_2c^2\cos\theta = Mc^2 \quad (\text{S.173})$$

e che, poiché il fotone ha massa nulla $E_i = p_i c$ e l'equazione diventa

$$2E_1E_2 - 2E_1E_2\cos\theta = Mc^2 \quad (\text{S.174})$$

da cui si ricava che

$$M = \frac{2E_1E_2(1 - \cos\theta)}{c^2}. \quad (\text{S.175})$$

Basta dunque conoscere l'energia e la direzione dei due fotoni per sapere quale particella ha dato origine al decadimento.

Esercizio 19.1

Il numero di particelle misurate con un ritardo t è proporzionale alla frequenza $f(t)$: $N(t) = Af(t)$. Nella formula (19.13) compaiono solo i rapporti tra i numeri di particelle $N(t)$ e $N(0)$ che dunque sono uguali ai rapporti $f(t)/f(0)$.

Possiamo dunque assumere $N(0) = 3.47$ e la vita media si può stimare dalla misura a $t = 1.00 \mu\text{s}$:

$$\tau_1 = -t \frac{1}{\log \frac{N(t_1)}{N(0)}} = -1.00 \frac{1}{\log \frac{2.42}{3.47}} \simeq 2.77 \mu\text{s}. \quad (\text{S.176})$$

Usando la misura a $t = 1.97 \mu\text{s}$ si trova $\tau_2 \simeq 1.94 \mu\text{s}$, mentre dalla misura a $t = 3.80 \mu\text{s}$ si ottiene $\tau_3 \simeq 2.72 \mu\text{s}$.

Il valore della vita media del muone si può quindi stimare come la media di questi, che vale $\langle \tau \rangle \simeq 2.48 \mu\text{s}$. L'errore da associare a questo numero si ricava valutando la varianza che vale $\sigma_\tau = 0.22 \mu\text{s}$ e dividendola per la radice del numero di misure: $\delta_\tau = \sigma_\tau / \sqrt{3} \simeq 0.16 \mu\text{s}$.

Il risultato si esprime come

$$\tau_\mu = 2.48 \pm 0.16 \mu\text{s}. \quad (\text{S.177})$$

La vita media del muone, come risulta da **misure** eseguite fino al 2012 è

$$\tau_\mu = 2.1969811 \pm 0.0000022 \mu\text{s}, \quad (\text{S.178})$$

valore che differisce da quello da noi trovato di circa $0.28 \mu\text{s}$ (meglio di due deviazioni standard).

Esercizio 20.1

Il pione carico **ha una vita media** pari a circa $\tau \simeq 10^{-8}$ s. Una volta prodotto si muove, nel sistema del laboratorio, con quantità di moto p che dipende dall'energia dei protoni usati.

Il decadimento è un processo stocastico perciò il pione può decadere, in linea di principio, in un tempo che va da 0 a ∞ . In media, circa 1/3 dei pioni sarà decaduto nel tempo di una vita media, data la forma della **legge del decadimento** e ricordando che $e \simeq 3$.

Lo spazio percorso in media in questo tempo vale

$$L \simeq v\tau = \frac{p}{m}\tau. \quad (\text{S.179})$$

Se $p \gg m$, $v \simeq c$ e quindi

$$L \simeq c\tau = 3 \times 10^8 \frac{\text{m}}{\text{s}} \times 10^{-8} \text{s} = 3 \text{ m}. \quad (\text{S.180})$$

In realtà, in questo caso, bisogna considerare che il tempo τ , nel sistema di riferimento del laboratorio, si dilata di un fattore

$$\gamma = \frac{1}{\sqrt{1 - \frac{v^2}{c^2}}} \quad (\text{S.181})$$

e quindi la distanza L può essere molto più lunga. Calcoliamo questo fattore sapendo che

$$E^2 = p^2 c^2 + m^2 c^4 = (\gamma^2 v^2 + c^2) m^2 c^2. \quad (\text{S.182})$$

Ma ora

$$\gamma^2 = \frac{1}{1 - \frac{v^2}{c^2}} = \frac{c^2}{c^2 - v^2} \quad (\text{S.183})$$

e sostituendo

$$E^2 = \left(\frac{c^2}{c^2 - v^2} v^2 + c^2 \right) m^2 c^2 = \frac{c^2}{c^2 - v^2} m^2 c^4 = \gamma^2 m^2 c^4 \quad (\text{S.184})$$

e pertanto

$$\gamma = \frac{E}{mc^2}, \quad (\text{S.185})$$

vale a dire γ rappresenta la frazione di energia in più rispetto alla massa a riposo della particella. Nel nostro caso abbiamo dunque

$$L = \gamma c\tau = \frac{E}{mc} \tau. \quad (\text{S.186})$$

Un pione di $300 \text{ MeV}/c$ di quantità di moto, avendo una massa pari a circa $140 \text{ MeV}/c^2$ ha un'energia pari a

$$E = \sqrt{(300 \text{ MeV}/c)^2 c^2 + (140 \text{ MeV}/c^2)^2 c^4} \simeq 330 \text{ MeV} \quad (\text{S.187})$$

e quindi il suo γ vale

$$\gamma = \frac{E}{mc^2} \simeq \frac{330 \text{ MeV}}{140 \frac{\text{MeV}}{c^2} c^2} \simeq 2.4. \quad (\text{S.188})$$

Lo spazio a disposizione per accelerare i pioni è di poco superiore ai 7 m ($2.4 \times 3 \text{ m}$).

Esercizio 20.2

Per permettere la produzione di una particella si devono conservare energia e quantità di moto. Nello stato iniziale abbiamo che l'energia iniziale E_{in} è pari alla somma dell'energia del pione e di quella associata alla massa a riposo m del [protone](#), cioè

$$E_{in} = E_{\pi} + mc^2. \quad (\text{S.189})$$

Quest'energia deve essere uguale a quella dello stato finale E , nel quale si trova una particella di massa M , che si muove con quantità di moto P , la cui energia è $E = \sqrt{P^2c^2 + M^2c^4}$. Inoltre si deve conservare la quantità di moto e quindi

$$p_{\pi} = P, \quad (\text{S.190})$$

perciò

$$E^2 = p_{\pi}^2c^2 + M^2c^4 = E_{\pi}^2 + m^2c^4 + 2mE_{\pi}c^2 \quad (\text{S.191})$$

e sostituendo l'energia del pione con la sua espressione $E_{\pi}^2 = p_{\pi}^2c^2 + m_{\pi}^2c^4$

$$p_{\pi}^2c^2 + M^2c^4 = p_{\pi}^2c^2 + m_{\pi}^2c^4 + mc^2 + 2mE_{\pi}c^2 \quad (\text{S.192})$$

da cui si ottiene che

$$E_{\pi} = \frac{M^2c^4 - m_{\pi}^2c^4 - mc^2}{2mc^2}. \quad (\text{S.193})$$

Quando dunque il pione assume quest'energia si può formare una particella di massa M . È in corrispondenza di questi casi che si osserva l'aumento della sezione d'urto. Usando l'equazione ([S.193](#)) si può quindi conoscere la massa della particella che si è formata conoscendo l'energia del pione in corrispondenza della quale la sezione d'urto presenta il valore massimo.

Esercizio 20.3

Per produrre le risonanze Δ con fasci di pioni si deve conservare, oltre alla carica elettrica, anche il

numero barionico. Le Δ sono barioni perciò, se nello stato finale c'è solamente una Δ , il numero barionico dello stato iniziale deve essere $B = +1$, che si può ottenere con un neutrone oppure un protone. Il segno del pione dipende quindi dalla risonanza che si vuole produrre. Per la Δ^- , ad esempio, la carica complessiva dello stato iniziale deve valere -1 in unità di carica del protone, perciò la reazione è

$$\pi^- + n \rightarrow \Delta^-. \quad (\text{S.194})$$

Analogamente, per produrre la Δ^0 si devono usare pioni negativi su protoni:

$$\pi^- + p \rightarrow \Delta^0. \quad (\text{S.195})$$

Le altre due reazioni sono

$$\pi^+ + n \rightarrow \Delta^+ \quad (\text{S.196})$$

e

$$\pi^+ + p \rightarrow \Delta^{++}. \quad (\text{S.197})$$

Esercizio 21.1

Per calcolare l'energia minima che dovrebbe avere un fascio di pioni per produrre un mesone K possiamo procedere analogamente a quanto fatto nella [soluzione](#) dell'Esercizio 18.1.

Consideriamo l'urto di un pione di massa m e con velocità v con un protone fermo nel sistema di riferimento del laboratorio. L'energia necessaria a produrre una risonanza con la massa del mesone K è almeno

$$E = m_Kc^2 \quad (\text{S.198})$$

dove m_K è la massa del K . Questa energia deve essere disponibile nel centro di massa del sistema pione-protone, in cui il protone è fermo e il pione si muove. Ricordando che la massa di una particella è un invariante relativistico possiamo affermare che

$$E^2 = E_{in}^2 - p_{in}^2c^2 \quad (\text{S.199})$$

dove E_{in} è l'energia complessiva dello stato iniziale composto del pione e del protone, mentre p_{in} la quantità di moto di questo stesso stato. Nello stato iniziale l'unica particella in moto è il pione perciò

$$p_{in} = p_{\pi} \quad (\text{S.200})$$

mentre all'energia contribuiscono sia il pione che il protone:

$$E_{in} = E_{\pi} + m_p c^2 = \sqrt{p_{\pi}^2 c^2 + m_{\pi}^2 c^4} + m_p c^2. \quad (\text{S.201})$$

Sostituendo l'espressione di E in termini della massa del K otteniamo che

$$m_K^2 c^4 = E_{\pi}^2 + m_p^2 c^4 + 2E_{\pi} m_p c^2 - p_{\pi}^2 c^2 \quad (\text{S.202})$$

e osservando che

$$E_{\pi}^2 - p_{\pi}^2 c^2 = m_{\pi}^2 c^4 \quad (\text{S.203})$$

si arriva all'espressione

$$m_K^2 c^4 = m_{\pi}^2 c^4 + m_p^2 c^4 + 2E_{\pi} m_p c^2 \quad (\text{S.204})$$

dalla quale si ricava l'energia che deve avere il pione:

$$E_{\pi} = \frac{m_K^2 - m_{\pi}^2 - m_p^2}{2m_p} c^2 \quad (\text{S.205})$$

Bibliografia

- [1] Donald Knuth. *The Art of Computer Programming*, volume 1–4. Addison–Wesley Professional, 1997.
- [2] Adriano Filipponi. *Introduzione alla fisica*. Zanichelli, 2005.
- [3] Agostino. *Le Confessioni*. Einaudi.
- [4] F. Parmeggiani, G. e Bonoli. Quirico filopanti: una singolare figura di astronomo nella bologna dell’ottocento. *Memorie della Società Astronomica Italiana*, 66(4):861–870, 1995.
- [5] Giulio Giorello. *Introduzione alla filosofia della scienza*. Bompiani, 2006.
- [6] F. Rabaud, M. e Moisy. Ship wakes: Kelvin or mach angle? *Phys. Rev. Lett.*, 110:214503, 2013.
- [7] Isaac Newton. *Principi matematici della filosofia naturale*, a cura di A. Pala. UTET, 1725.
- [8] Augustin Fresnel. *The wave theory of light*. Cincinnati American Book Company, 1900.
- [9] Thomas Young. *Experiments and calculations relative to physical optics*. *Phil. Trans. R. Soc. Lond.*, 94:1,16, 1803.
- [10] Giovanni Organtini. *Navigatori GPS e teoria della relatività*. Accastampato, 9, 2012.
- [11] R. Ariew. *G. W. Leibniz and Samuel Clarke: Correspondence*. Hackett Publishing Co. Inc. Indianapolis/Cambridge, 2000.
- [12] Max Planck. On the law of distribution of energy in the normal spectrum. *Annalen der Physik*, 4:553, 1901.
- [13] Albert Einstein. *Über einen die Erzeugung und Verwandlung des Lichtes betreffenden heuristischen Gesichtspunkt*. *Annalen der Physik*, 17(6):132, 1905.
- [14] O. Gerlach, W. e Stern. Das magnetische moment des silberatoms. *Zeitschrift für Physik*, 9(1):353.
- [15] Domenico Pacini. *La radiazione penetrante alla superficie ed in seno alle acque*. *Il Nuovo Cimento, Series VI*, 3(93D100), 1912.
- [16] Viktor F. Hess. *über beobachtungen der durchdringenden strahlung bei sieben freiballonfahrten*. *Physikalische Zeitschrift*, 13:1084, 1091, 1912.
- [17] R. A. Millikan. *High frequency rays of cosmic origin*. 12(1):48, 55, 1925.
- [18] A. M. Hillas. *Cosmic Rays: Recent Progress and some Current Questions*. *arXiv:astro-ph/0607109 v2*, 2006.
- [19] Enrico Fermi. *On the Origin of the Cosmic Radiation*. *Physical Review*, 75:1169, 1174, 1949.
- [20] D. V. Skobelzyn. *A New Type of Very Fast Beta Rays*. *Z. Phys.*, 54:686, 1929.
- [21] L. Brown. *The idea of the neutrino*. *Phys. Today*, 31(9), 1978.
- [22] D. Anderson. *The Positive Electron*. *Physical Review*, 43(6):491, 494, 1933.

- [23] Paul Kunze. **Untersuzhung der Ultrastrahlung in der Wilsonkammer.** *Zeitschrift fur Physik*, 83(1-2):1, 18, 1933.
- [24] G. P. S Blackett, P. M. S. e Occhialini. **Some photographs of the tracks of penetrating radiation.** *Proceedings of the Royal Society of London. Series A, Containing Papers of a Mathematical and Physical Character*, 139, 1933.
- [25] S. H. Anderson, C. D. e Neddermeyer. **Note on the nature of cosmic-ray particles.** *Phys. Rev.*, 50:263, 1937.
- [26] J. C. e Stevenson E. C. Street. **New Evidence for the Existence of a Particle of Mass Intermediate Between the Proton and Electron.** *Phys. Rev.*, 52:1003, 1004, 1937.
- [27] C. C. Rochester, G. D. e Butler. **Evidence for the existence of new unstable elementary particles.** *Nature*, 160:855, 857, 1947.
- [28] V. D. e Biswas S. Hopper. **Evidence Concerning the Existence of the New Unstable Elementary Neutral Particle.** *Phys. Rev.*, 80:1099, 1950.
- [29] Oreste Conversi, Marcello e Piccioni. **Misura diretta della vita media dei mesoni frenati.** *Il Nuovo Cimento*, 2(1):40, 70, 1944.
- [30] Murray Gell-Mann. **The Eightfold Way: A Theory of Strong Interaction Symmetry.** *Synchrotron Laboratory Report*, (CTSL-20), 1961.
- [31] S. Okubo. **Note on unitary symmetry in strong interactions.** *Progress in theoretical physics*, 27(5):949, 1962.
- [32] Giovanni Organtini. **Unveiling the Higgs mechanism to students.** *Eur. J. Phys.*, 33:1397, 2012.

Storia delle edizioni

Questa è la **terza edizione** di questo volume. La seconda edizione includeva correzione di errori e revisione di alcuni capitoli dopo aver ricevuto alcuni commenti sulla prima edizione, ed è stata pubblicata nel settembre 2014.

In questa edizione sono stati aggiunti i capitoli introduttivi, quelli sulle onde e quello sull'entropia.

Questo testo è anche il risultato del contributo di coloro che hanno segnalato errori o imprecisioni, o che hanno suggerito modi alternativi o più comprensibili di spiegare alcuni concetti. Di seguito sono indicati i contributori e gli interventi segnalati, in ordine cronologico.

1. Carlo Sciò: che ha segnalato alcuni refusi.
2. Anonimo: al Paragrafo "Semiconduttori", nei cristalli di tipo p gli elettroni si trovano nella banda di valenza e non in quella di conduzione.
3. Ciro Chiaiese: segnalati errori nel Paragrafo [12.2](#); alcune equazioni erano sbagliate.
4. Claudia Pinzari: suggerita una riformulazione del metodo per giungere alle trasformazioni di Lorentz.

Abbiamo inoltre apportato alcune modifiche in paragrafi che giudicavamo poco chiari o che sono risultati tali dopo aver discusso con alcuni lettori. In particolare in questa versione è stata abbastanza rivisto il paragrafo che illustra l'emissione del corpo nero e quello nel quale si illustrano gli effetti relativistici sulla massa classica.

Indice analitico

Čerenkov, effetto, 88

Abbe, limite di, 104

acceleratore, 143, 195, 196, 203, 210, 222, 233

accuratezza, 22

acutezza visiva, 104

adimensionale, 20

adrone, 192

α , raggi, 177, 180

ampiezza

di un'onda, 76

amplificatore, 171

Anderson, Carl, 186, 187

angolo critico, 69

annichilazione, 230

anno luce, 127

antenna, 153

$\bar{K}$, 213, 214, 216

antimateria, 186, 230

$\bar{\nu}$, 186

$\bar{p}$, 186

antiprotone, 186

apparente, forza, 146

arcobaleno, 65

armonica, 74, 83

asse ottico, 61

atomo, 219

autovelox, 86

Avogadro, Numero di, 124

Avogadro, numero di, 169, 219

banda, 102, 169

di conduzione, 170

di valenza, 170

barione, 191, 192, 213, 214, 216

barionico, numero, 191, 196, 213, 221

barn, 197

base, 171

battimenti, 81

Bayes, Thomas, 27

beauty, quark, 221

berillio, 196

β , decadimento, 184, 185

β , raggi, 177, 179, 184, 185

^{210}Bi , 184

bianca

particella, 218

bilancia, 40

bin, 30

binocolo, 72

bischero, 82

Bismuto 210, 184

Biswas, S, 189

Blackett, Patrick, 187

Blu-Ray, 103

Boltzmann, costante di, 119

bosone, 193

di Higgs, 143, 233

vettore, 233, 240

W, 231

Z, 231

bottom, quark, 221

Bracciano, Lago di, 178

branching ratio, 199

bureau international de poids et mesures, 18

Butler, Clifford, 188

calibrazione, 22

calore

latente, 114

calore latente, 57

calore specifico, 53

calorico, 51

camera

a nebbia, 183, 188

a scintilla, 180

- di Wilson, 183
- campo
 - magnetico, 183, 186, 188, 233
 - terrestre, 179
- canale
 - di decadimento, 199
- capacità termica, 53
- Carnot, ciclo di, 110
- Carnot, macchina di, 109
- Carnot, Sadi, 110
- CD, 102
- Celsius, grado, 48
- centrifuga, forza, 147
- Čerenkov, 90
- charm, quark, 221
- Čerenkov, 90
- chi quadro, 42
- χ^2 , 42
- chitarra, 82
- ciclo di Carnot, 110
- cladding, 72
- clarinetto, 102
- Clausius, integrale di, 113
- Clausius, Rudolf, 113
- ^{60}Co , 184
- cobalto, 41
- Cobalto 60, 184
- Collegno, Lo smemorato di, 87
- collettore, 171
- collider, 143, 144
- collisore, 143, 144
- colore, 65, 217, 218
- composizione delle velocità, 130
- Compton, effetto, 157
- Compton, lunghezza d'onda, 158
- conducibilità termica, 54
- conduzione, 54
- confinamento, 219
- conservazione
 - del numero barionico, 191, 213, 221, 257
 - del numero leptonico, 191, 221, 257
 - dell'energia, 185, 192, 201, 224, 257
 - della carica, 184, 187, 191, 257
 - della massa, 41
 - della quantità di moto, 94, 181, 185, 257
 - della stranezza, 210
- contrazione della lunghezza, 129
- convezione, 55
- Copernico, Niccolò, 16
- corda, 82
- core, 72
- corpo nero, 153
- corrente, 169
- cosmici, provenienza dei raggi, 181
- cosmici, raggi, 128, 177, 179, 180, 184, 187, 189, 195, 209, 222
- costante
 - di Boltzmann, 119
 - di Planck, 193, 246
 - di Stefan–Boltzmann, 56
- Creative Commons, 3
- cristallo drogato, 170
- Ξ , 211, 213, 215
- Ξ^* , 213
- debole, forza, 198
- debole, interazione, 198
- debole, interazione o forza, 185, 231
- decadimento, 197, 198
 - β , 184
 - del muone, 187
- decadimento β , 185
- decadimento radioattivo, 41
- decupletto di barioni, 214
- deferente, 17
- definizione operativa, 197
- Δ , 205, 213, 215, 218
- derivata, grandezza fisica, 20
- deviazione standard, 31, 123
- diagramma di Feynman, 227
- diapason, 82
- diffrazione, 93, 94
 - dei suoni, 101
 - di onde luminose, 102
 - reticolo di, 100
- dilatazione del tempo, 127
- dilatazione termica, 48
- dimensione fisica, 20
- dimensioni fisiche, 20
- diodo, 170

- Dirac, Paul, 186
disco di Newton, 65
disordine, 122
dispersione, 65, 68
distribuzione uniforme, 32, 122
Doppler, ecografia, 86
Doppler, effetto, 85
down, quark, 215
drogato, cristallo, 170
Duchamp, Marcel, 15
DVD, 103
- eco, 93
ecodoppler, 86
ecografia Doppler, 86
effetto Čerenkov, 88
effetto Compton, 157
effetto Doppler, 85
effetto est-ovest, 179
effetto fotoelettrico, 155
Einstein, Albert, 14, 125, 156, 219
elettrone, 180, 183, 184
elettroscopio, 177
eliocentrica, teoria, 16
emettitore, 171
energia, 224
 di soglia, 196, 258
 meccanica, 224
energia a riposo, 135
energia-impulso, quadrivettore, 136
energia-impulso, tensore, 150
entropia, 112
 dell'Universo, 113
epiciclo, 17
equazione
 differenziale, 245
 a variabili separabili, 246
equazione degli specchi, 71
equazione della lente, 71
equazione di Schrödinger, 172, 174
equazione dimensionale, 20
equazioni di Maxwell, 127, 163
equivalenza, principio di, 146
errore, 28
 propagazione dell', 255
 relativo, 255
errore relativo, 66
errore sistematico, 32
esclusione
 principio di, 166
esclusione, principio di, 165
esperimento, 15
est-ovest, effetto, 179
estimatore, 29
 η , 213, 214, 216
etere, 127
euristica positiva, 17
evento, 132, 139
- fase, 80
fasi, spazio delle, 200
fattore relativistico di Lorentz, 184, 257
fenditura, 97
Fermi, Enrico, 181, 184
Fermi, Modello di, 181
fermione, 193
Feynman, diagramma di, 227
FFT, 102
fibra ottica, 59, 72
Filopanti, Quirico, 39
Finnegans Wake, 213
fit lineare, 43
Fiumicino, 95
fondamentale, armonica, 83
fondamentale, grandezza fisica, 18
fondo scala, 22
forte, forza, 192, 198
forte, interazione, 192, 198
forte, interazione o forza, 192
forza
 centrifuga, 147
 debole, 198, 223
 di gravità, 223
 di Lorentz, 135, 179, 183, 256
 di Van der Waals, 219
 elettromagnetica, 223
 forte, 192, 198, 223
 gravitazionale, 223
forza apparente, 146
forza debole, 185, 231

- forza forte, 192
fotoelettrico, effetto, 155
fotone, 156, 157, 193, 229
fotonica, 104
Fourier, teorema di, 74
Fourier, trasformata di, 102
frequenza, 25, 122
 di un'onda, 76
fronte d'onda, 77
funzione
 di stato, 109, 112
funzione d'onda, 173
fuoco, 61
fuoco (di una lente), 70
fusi orari, 39
- Galileo, trasformazioni di, 127, 130
 γ di Lorentz, 184, 257
 γ , raggi, 180
Gauss (u.m.), 256
Gauss, curva di, 28
gaussiana, 28
gaussiana, funzione, 123
Gell-Mann, Murray, 213
gemelli, paradosso dei, 130, 152
generale, relatività, 130, 145
geocentrica, teoria, 16
Gerlach, Walther, 164
ginocchio, 180
gluone, 231
gnomone, 39
GPS, 128, 151, 249
grancassa, 102
grandezza fisica, 15
 derivata, 20, 246
 fondamentale, 18, 246
grandi numeri, legge dei, 26
gravità, 223
- Hahnemann, Samuel, 8
Hertz, Heinrich, 155
Hess, Viktor, 177, 179
Higgs
 bosone di, 233
 meccanismo di, 233
- Higgs, bosone di, 143, 233
Higgs, Peter, 233
Hopper, V, 189
Huygens, principio di, 95
- immagine reale, 63
immagine virtuale, 63
impuslo, 134
incertezza, 28
indeterminazione, 28
indeterminazione, principio di, 158
indice di rifrazione, 68, 78
inerziale, sistema, 130, 145, 146
integrale
 di Clausius, 113
intensità, 100
interazione
 debole, 198, 201, 223
 elettromagnetica, 223
 forte, 192, 198, 223
 gravitazionale, 223
interazione debole, 185, 231
interazione forte, 192
interferenza, 97
interferenze, 79, 98
invariante, 132
invariante di Lorentz, 132
ionizzazione, 177
irraggiamento, 56
isobara, trasformazione, 116
isocora, trasformazione, 116, 117
isolante, 169
isoterma, trasformazione, 116, 117
istogramma, 25, 30
- jamendo, 3
Joyce, James, 213
- K , 211, 216
 K , 213
 K , 213
Kelvin, 89
Knuth, Donald, 3
Kolhörster, Werner, 179
Kunze, Paul, 187

- lacuna, 170
- Lago di Bracciano, 178
- Lakatos, Imre, 17
- Λ , 209–211, 213, 214
- larghezza
 - di una particella, 160
- laser, 59, 102
- L^AT_EX, 3
- Lavoisier, Antoine, 41
- Lavoisier, Legge di, 41
- legge
 - dei grandi numeri, 26
- Legge di Newton, 164
- legge di Snell, 68
- legge fisica, 16
- lente, 69
 - convergente, 70
 - d'ingrandimento, 70
- lente d'ingrandimento, 70
- lente sottile, 69
- leptone, 191, 198
 - τ , 221
- leptonico, numero, 191, 221
- limite centrale, teorema, 123
- limite di Abbe, 104
- Livorno, 178
- Lo smemorato di Collegno, 87
- Lorentz, fattore relativistico, 184, 257
- Lorentz, Forza di, 256
- Lorentz, forza di, 135, 179, 183
- Lorentz, invariante di, 132
- Lorentz, trasformazioni di, 127, 139
- luce, 201
 - luce, diffrazione della, 102
 - luce, natura della, 93
 - luce, velocità della, 136
- lunghezza d'interazione, 198
- lunghezza d'onda, 76
- lunghezza d'onda Compton, 158
- lunghezza, contrazione della, 129
- macchina, 107
 - di Carnot, 109
- macrostato, 120
- magnetico
 - campo, 233
- Malerba, Luigi, 125, 157
- mano destra, regola della, 179
- mantissa, 23
- massa, 18, 39
 - invariante, 204
- massa a riposo, 135
- massa invariante, 196, 206
- massa nulla, 136
- materializzazione, 195, 230
- Maxwell, equazioni di, 127, 163
- meccanismo di Higgs, 233
- media, 29
- media pesata, 37
- Mendelev, Dmitrii, 166, 213
- Mercurio, 17
- meridiana, 39
- mesone, 192, 213, 214
- metodo sperimentale, 7
- Michelson, Albert Abraham, 125
- micrometro, 18
- micron, 18
- microscopio, 71
- microscopio elettronico, 160
- microstato, 120
- Millikan, Robert, 179
- Minkowski, spazio di, 140
- Minkowski, spazio di, 132
- mirascopio, 65
- misura, 15
 - indiretta, 21
- Modello di Fermi, 181
- modulo, 31
- momento angolare, 161
- Morley, Edward, 125
- moto
 - retrogrado, 17
- moto browniano, 219
- moto perpetuo, 112
- MPEG Streamclip, 3
- μ , 188
- μ , 187, 188, 191, 199–201, 205, 221
- muone, 128, 159, 187, 188, 191, 199–201, 205, 221
- Muster Mark, 213

- nanometro, 18
- Neddermeyer, Seth, 187
- neutrino, 137, 184, 186
- Newton, disco di, 65
- Newton, legge di, 164
- ^{60}Ni , 184
- Nichel 60, 184
- nodo, 83
- normalizzazione, 25
- notazione scientifica, 23
- ν , 186
- numero
 - barionico, 191, 196, 213, 221
 - leptonico, 191, 221
- numero d'onda, 80, 159
- Numero di Avogadro, 124
- numero di Avogadro, 169, 219
- obiettivo, 71
- Occhialini, Giuseppe, 187
- oculare, 71
- Okubo, Susumo, 213
- Ω , 211, 215, 216
- Ω^- , 213
- omeopatia, 8
- onda, 73
 - stazionaria, 83
 - trasversale, 73
- onda elettromagnetica, 201
- Open Source, 3
- ordine di grandezza, 23
- ottetto
 - di barioni, 214, 216
 - di mesoni, 214, 216
- ottica, 59
- ottometrica, tavola, 104
- Pacini, Domenico, 178, 179
- pallone aerostatico, 179
- paradosso
 - dei gemelli, 130, 152
- Parmenide, 14
- particella bianca, 218
- particelle strane, 209, 210
- Pauli, Principio di, 217, 218
- Pauli, principio di, 165, 166
- Pauli, Wolfgang, 166, 217, 218
- pentaprisma, 72
- Penzias, Arno Allan, 34
- perielio, precessione del, 17
- Periodica, Tavola, 213
- periodica, tavola, 166
- periodo
 - di un'onda, 76
- Perrin, Jean Baptiste, 219
- peso, 18, 39
- Physics Gizmo, 24
- π , 187, 188, 192, 196, 198, 200, 201, 203–205, 209–211, 213, 214, 216, 260
- pinta, 21
- piolo, 82
- pione, 187, 188, 192, 196, 198, 200, 201, 203–205, 209–211, 213, 214, 216, 260
- pirolo, 82
- pixel, 171
- Planck, costante di, 155, 193, 246
- Planck, Max, 154
- ^{210}Po , 184
- Polonio 210, 184
- ponticello, 82
- portata, 22
- positrone, 186, 230
- potenziale, 225
- potere risolutivo, 104
- precessione
 - del perielio, 17
- precisione, 22
- primari, raggi cosmici, 180, 181, 195
- principio
 - di equivalenza, 146
 - di esclusione, 166, 217, 218
 - di Pauli, 217, 218
- principio d'indeterminazione, 158
- principio di esclusione, 165
- principio di Huygens, 95
- principio di Pauli, 165
- principio zero della termodinamica, 46
- prisma, 90
- probabilità, 26
- propagazione

- degli errori, 255
- prossimità, sensore di, 24
- prostaferesi, 81, 83
- protone, 137, 180, 184, 185, 187, 191, 192, 196, 201, 210, 211, 218, 221, 257, 258, 260
- provenienza dei raggi cosmici, 181
- pulsazione, 77
- quadratura, 36
- quadrimpulso, 136
- quadrivelocità, 133
- quadrivettore, 131, 132, 139
- quadrivettore energia-impulso, 136
- quantità di moto
 - conservazione della, 94
- quantizzazione dell'energia, 162
- quanto, 154
- quark, 144, 213, 215
 - beauty, 221
 - bottom, 221
 - charm, 221
 - down, 215
 - strange, 215
 - top, 221
 - up, 215
- Rabi, Isidor, 187
- radiazione
 - α , 177
 - β , 177, 179, 184, 185
- radiazione cosmica di fondo, 34
- radioattività, 41, 177
- raggi
 - γ cosmici, 128
 - α , 177, 180
 - β , 177, 179, 184, 185
 - cosmici, 177, 179, 180, 184, 187, 189, 195, 209, 222
 - primari, 180, 181, 195, 196
 - secondari, 181
 - solari, 181
 - γ , 180
 - X, 181
- raggi cosmici, sorgenti, 181
- raggi cosmici: provenienza, 181
- raggi X, 104
- rapporto di diramazione, 199
- rapporto incrementale, 245
- reale, immagine, 63
- red shift, 88, 90
- reflex, 72
- regola
 - della mano destra, 179
- regressione lineare, 42, 43
- relatività
 - generale, 145
 - ristretta, 125
- relatività generale, 130
- relatività speciale, 125
- relatività, teoria della, 185, 195
- rendimento, 108
- reticolo di diffrazione, 100
- retrogrado, moto, 17
- reversibile, trasformazione, 109
- Riemann, tensore di, 150
- riflessione, 59, 77, 93, 94
- riflessione totale, 69
- riflettanza, 60
- riflettività, 60
- rifrazione, 66, 77, 93, 94
- rifrazione, indice di, 68, 78
- riga
 - di assorbimento, 91
- riga, in uno spettro, 162
- risolutivo, potere, 104
- risoluzione, 104
- risonanza, 207
- risonanze, 203, 204
- ristretta, relatività, 125
- Rochester, George, 188
- Rossi, Bruno, 179
- rottura della simmetria, 241
- sagitta, 255
- scarica
 - degli elettroscopi, 177
- Schlieren flow, 94
- Schrödinger, equazione di, 172, 174
- Schrödinger, Erwin, 172
- sciacquone, 171

- scintilla, camera a, 180
- sec:interpretazione-misure, 28
- secondari, raggi cosmici, 181
- secondo principio della termodinamica, 112
- semiconduttori, 169
- semidispersione massima, 29
- sensibilità, 22
- sezione d'urto, 197, 203
 - elettromagnetica, 198
 - forte, 198
- SI, 197, 246
- Σ , 213, 214
- Σ^* , 213, 215
- Sistema Internazionale, 197, 246
- sistema internazionale, 18
- sistematico, errore, 32
- Skobeltsyn, Dmitri, 183
- smartphone, 102
- Snell, legge di, 68
- Snellius, Willebrord, 68
- soglia, energia di, 196, 258
- solari, raggi cosmici, 181
- sorgente
 - di raggi cosmici, 181
- spazio
 - delle fasi, 201
- spazio curvo, 149
- spazio delle fasi, 200
- spazio piatto, 149
- specchio sferico, 60
- speciale, relatività, 125
- sperimentale, metodo, 7
- spettro, 91, 102, 162
 - dei raggi cosmici, 180
 - del decadimento β , 185
- spin, 164, 193
- spostamento
 - verso il rosso, 88
- stazionaria, onda, 82, 83
- Stazione zero, 125
- Stefan-Boltzmann, costante di, 56
- Stern, Otto, 164
- Stern-Gerlach, esperimento di, 164
- stocastico, 199
- strane, particelle, 209, 210
- stranezza, 209–211, 214, 231
- strange, 215
- strumento
 - graduato, 21
 - tarato, 22
- suono, 73
- suono, diffrazione del, 101
- Super-Kamiokande, 90
- supernovae, 181
- taratura, 22
- τ , leptone, 221
- tavola ottometrica, 104
- Tavola Periodica, 213
- tavola periodica, 166
- telescopio, 71
- temperatura, 46
- tempo, 18
- tempo di decadimento, 45
- tempo di dimezzamento, 46
- tempo proprio, 128
- tempo, dilatazione del, 127
- tensore, 145
 - di Riemann, 150
 - energia-impulso, 150
 - metrico, 146
- teorema
 - del limite centrale, 123
 - di Fourier, 74
- teoria, 16
 - della relatività, 185, 195
 - eliocentrica, 16
 - geocentrica, 16
- termodinamica, principio zero, 46
- termodinamica, secondo principio, 112
- termometro, 48
- timpano, 74
- Tolomeo, Claudio, 16
- top, quark, 221
- Totò, 87
- transistor, 171
- trasformata di Fourier, 102
- trasformazione
 - isobara, 116
 - isocora, 116, 117

isoterma, [116](#), [117](#)
reversibile, [109](#)
trasformazioni di Galileo, [127](#), [130](#)
trasformazioni di Lorentz, [127](#)
trasversale, onda, [73](#)

uniforme, distribuzione, [122](#)
unità naturali, [130](#), [193](#), [246](#)
Universo, entropia dell', [113](#)
up, quark, [215](#)
Urano, [17](#)

valenza, [166](#)
valor medio, [29](#)
Van der Waals, forza di, [219](#)
Van der Waals, Johannes, [219](#)
varianza, [31](#)
 del campione, [31](#)
 della popolazione, [31](#)
velocità, [76](#)
 della luce, [136](#)
velocità della luce, [246](#)
velocità, composizione delle, [130](#)
violino, [82](#)
virtuale, immagine, [63](#)
vita media, [198](#)

W, [231](#)
Wilson, camera di, [183](#)
Wilson, Charles, [183](#)
Wilson, Robert Woodrow, [34](#)

X, raggi, [104](#), [181](#)
 Ξ , [211](#), [213](#), [215](#)
 Ξ^* , [213](#)

Z, [231](#)

